

VISIT...

LANZAROTE
Caliente.COM

CHAPTER

1

Introduction to Molecular Genetics and Genomics



CHAPTER OUTLINE

- 1.1 **DNA: The Genetic Material**
Experimental Proof of the Genetic Function of DNA
Genetic Role of DNA in Bacteriophage
- 1.2 **DNA Structure: The Double Helix**
- 1.3 **An Overview of DNA Replication**
- 1.4 **Genes and Proteins**
Inborn Errors of Metabolism as a Cause of Hereditary Disease
Mutant Genes and Defective Proteins
- 1.5 **Gene Expression: The Central Dogma**
Transcription
Translation
The Genetic Code
- 1.6 **Mutation**
Protein Folding and Stability
- 1.7 **Genes and Environment**
- 1.8 **Evolution: From Genes to Genomes, from Proteins to Proteomes**
The Molecular Unity of Life
Natural Selection and Diversity

PRINCIPLES

- Inherited traits are affected by genes.
- Genes are composed of the chemical deoxyribonucleic acid (DNA).
- DNA replicates to form copies of itself that are identical (except for rare mutations).
- DNA contains a genetic code specifying what types of enzymes and other proteins are made in cells.
- DNA occasionally mutates, and the mutant forms specify altered proteins that have reduced activity or stability.
- A mutant enzyme is an “inborn error of metabolism” that blocks one step in a biochemical pathway for the metabolism of small molecules.
- Traits are affected by environment as well as by genes.
- Organisms change genetically through generations in the process of biological evolution.

CONNECTIONS

Shear Madness

Alfred D. Hershey and Martha Chase 1952
Independent Functions of Viral Protein and Nucleic Acid in Growth of Bacteriophage

The Black Urine Disease

Archibald E. Garrod 1908
Inborn Errors of Metabolism

Each species of living organism has a unique set of inherited characteristics that makes it different from other species. Each species has its own developmental plan—often described as a sort of “blueprint” for building the organism—which is encoded in the DNA molecules present in its cells. This developmental plan determines the characteristics that are inherited. Because organisms in the same species share the same developmental plan, organisms that are members of the same species usually resemble one another, although some notable exceptions usually are differences between males and females. For example, it is easy to distinguish a human being from a chimpanzee or a gorilla. A human being habitually stands upright and has long legs, relatively little body hair, a large brain, and a flat face with a prominent nose, jutting chin, distinct lips, and small teeth. All of these traits are inherited—part of our developmental plan—and help set us apart as *Homo sapiens*.

But human beings are by no means identical. Many traits, or observable characteristics, differ from one person to another. There is a great deal of variation in hair color, eye color, skin color, height, weight, personality traits, and other characteristics. Some human traits are transmitted biologically, others culturally. The color of our eyes results from biological inheritance, but the native language we learned as a child results from cultural inheritance. Many traits are influenced jointly by biological inheritance and environmental factors. For example, weight is determined in part by inheritance but also in part by environment: how much food we eat, its nutritional content, our exercise regimen, and so forth. **Genetics** is the study of biologically inherited traits, including traits that are influenced in part by the environment.

The fundamental concept of genetics is that:

Inherited traits are determined by the elements of heredity that are transmitted from parents to offspring in reproduction; these elements of heredity are called **genes**.

The existence of genes and the rules governing their transmission from generation to generation were first articulated by Gregor Mendel in 1866 (Chapter 3). Mendel’s formulation of inheritance was in

terms of the abstract rules by which hereditary elements (he called them “factors”) are transmitted from parents to offspring. His objects of study were garden peas, with variable traits like pea color and plant height. At one time genetics could be studied only through the progeny produced from matings. Genetic differences between species were impossible to define, because organisms of different species usually do not mate, or they produce hybrid progeny that die or are sterile. This approach to the study of genetics is often referred to as *classical* genetics, or organismic or morphological genetics. Given the advances of *molecular*, or modern, genetics, it is possible to study differences between species through the comparison and analysis of the DNA itself. There is no fundamental distinction between classical and molecular genetics. They are different and complementary ways of studying the same thing: the function of the genetic material. In this book we include many examples showing how molecular and classical genetics can be used in combination to enhance the power of genetic analysis.

The foundation of genetics as a molecular science dates back to 1869, just three years after Mendel reported his experiments. It was in 1869 that Friedrich Miescher discovered a new type of weak acid, abundant in the nuclei of white blood cells. Miescher’s weak acid turned out to be the chemical substance we now call **DNA (deoxyribonucleic acid)**. For many years the biological function of DNA was unknown, and no role in heredity was ascribed to it. This first section shows how DNA was eventually isolated and identified as the material that genes are made of.

1.1 DNA: The Genetic Material

That the cell nucleus plays a key role in inheritance was recognized in the 1870s by the observation that the nuclei of male and female reproductive cells undergo fusion in the process of fertilization. Soon thereafter, **chromosomes** were first observed inside the nucleus as thread-like objects that become visible in the light microscope when the cell is stained with certain dyes. Chromosomes were found to exhibit a characteristic “splitting” behavior in which each daughter cell formed by cell division

receives an identical complement of chromosomes (Chapter 4). Further evidence for the importance of chromosomes was provided by the observation that, whereas the number of chromosomes in each cell may differ among biological species, the number of chromosomes is nearly always constant within the cells of any particular species. These features of chromosomes were well understood by about 1900, and they made it seem likely that chromosomes were the carriers of the genes.

By the 1920s, several lines of indirect evidence began to suggest a close relationship between chromosomes and DNA. Microscopic studies with special stains showed that DNA is present in chromosomes. Chromosomes also contain various types of proteins, but the amount and kinds of chromosomal proteins differ greatly from one cell type to another, whereas the amount of DNA per cell is constant. Furthermore, nearly all of the DNA present in cells of higher organisms is present in the chromosomes. These arguments for DNA as the genetic material were unconvincing, however, because crude chemical analyses had suggested (erroneously, as it turned out) that DNA lacks the chemical diversity needed in a genetic substance. The favored candidate for the genetic material was protein, because proteins were known to be an exceedingly diverse collection of molecules. Proteins therefore became widely accepted

as the genetic material, and DNA was assumed to function merely as the structural framework of the chromosomes. The experiments described below finally demonstrated that DNA is the genetic material.

Experimental Proof of the Genetic Function of DNA

An important first step was taken by Frederick Griffith in 1928 when he demonstrated that a physical trait can be passed from one cell to another. He was working with two strains of the bacterium *Streptococcus pneumoniae* identified as S and R. When a bacterial cell is grown on solid medium, it undergoes repeated cell divisions to form a visible clump of cells called a **colony**. The S type of *S. pneumoniae* synthesizes a gelatinous capsule composed of complex carbohydrate (polysaccharide). The enveloping capsule makes each colony large and gives it a glistening or smooth (S) appearance. This capsule also enables the bacterium to cause pneumonia by protecting it from the defense mechanisms of an infected animal. The R strains of *S. pneumoniae* are unable to synthesize the capsular polysaccharide; they form small colonies that have a rough (R) surface (**Figure 1.1**). This strain of the bacterium does not cause pneumonia, because without the capsule the bacteria are inactivated by the immune system of the host. Both types of bacteria

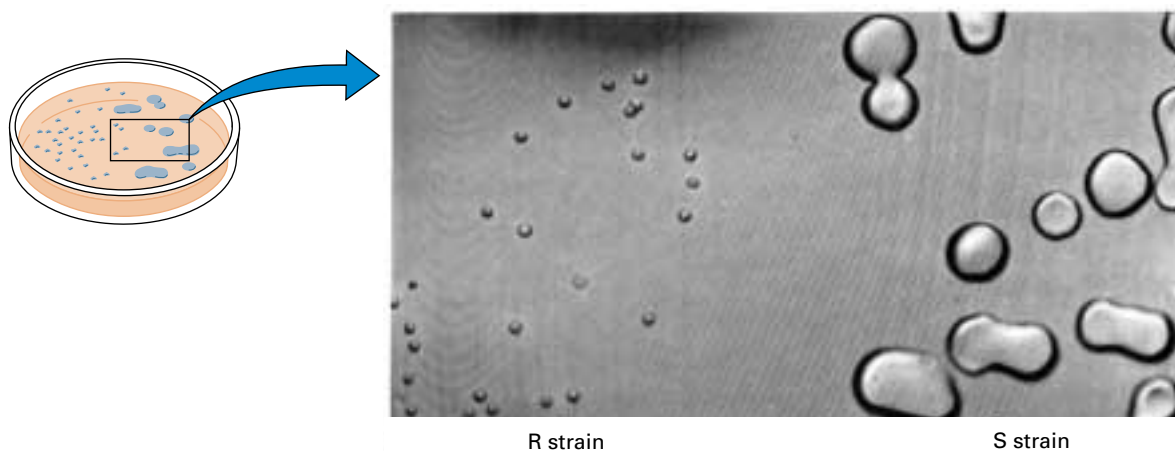


Figure 1.1 Colonies of rough (R, the small colonies) and smooth (S, the large colonies) strains of *Streptococcus pneumoniae*. The S colonies are larger because of the gelatinous capsule on the S cells. [Photograph from O. T. Avery, C. M. MacLeod, and M. McCarty. Reproduced from the *Journal of Experimental Medicine*, 1944, vol. 79, p. 137 by copyright permission of The Rockefeller University Press.]

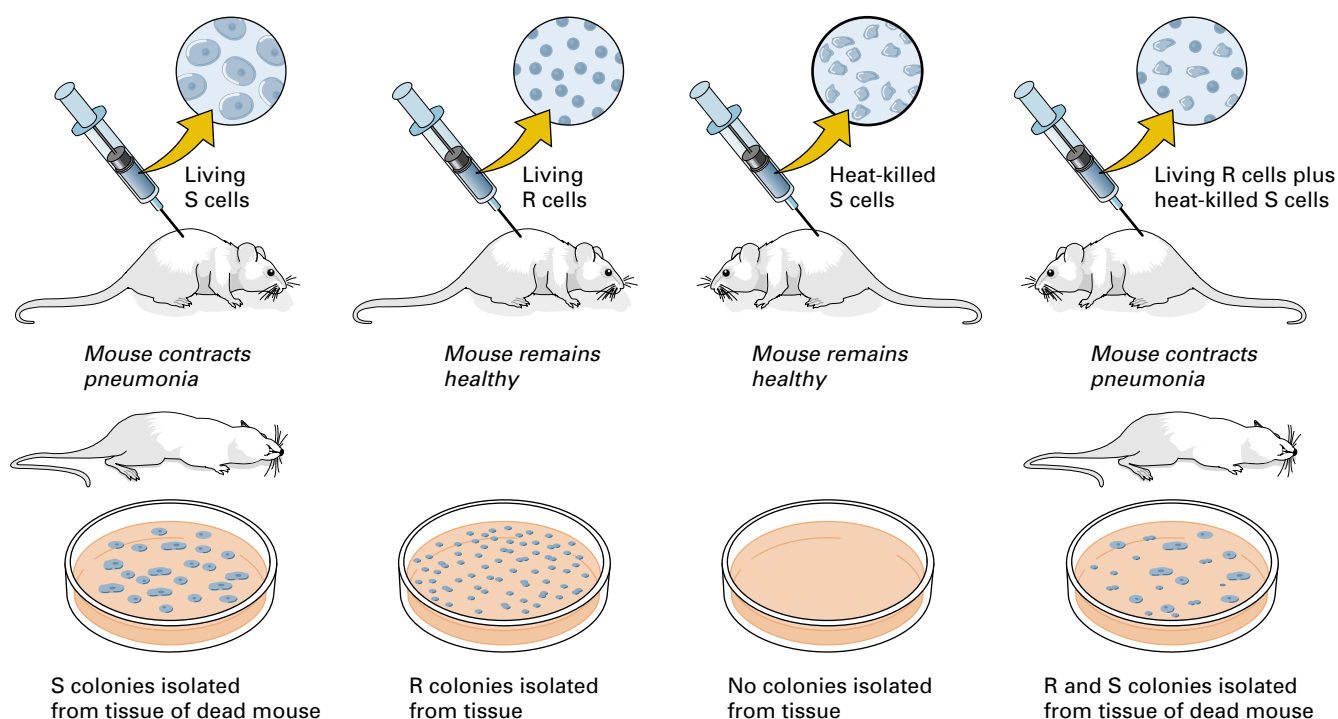


Figure 1.2 The Griffith's experiment demonstrating bacterial transformation. A mouse remains healthy if injected with either the nonvirulent R strain of *S. pneumoniae* or heat-killed cell fragments of the usually virulent S strain. R cells in the presence of heat-killed S cells are transformed into the virulent S strain, causing pneumonia in the mouse.

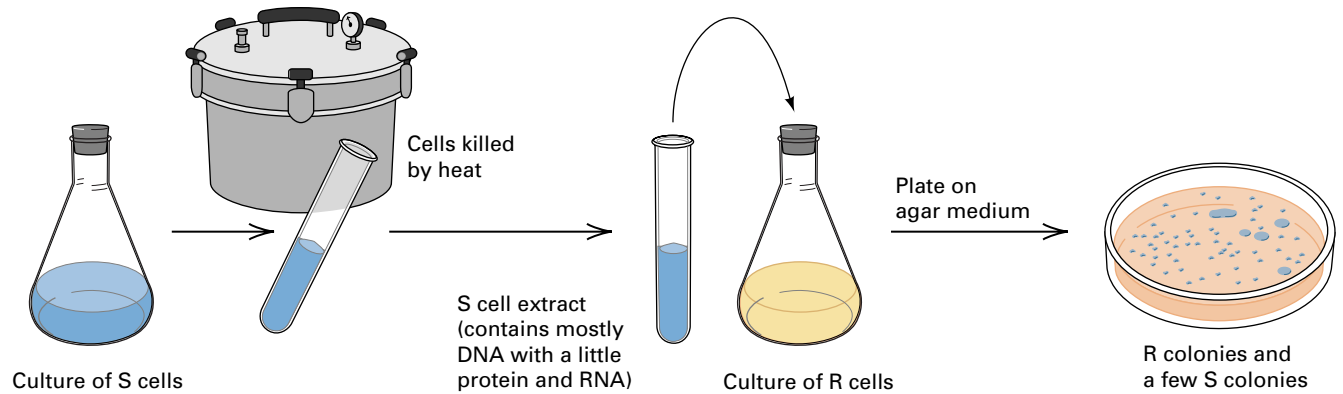
"breed true" in the sense that the progeny formed by cell division have the capsular type of the parent, either S or R.

Mice injected with living S cells get pneumonia. Mice injected either with living R cells or with heat-killed S cells remain healthy. Here is Griffith's critical finding: mice injected with a *mixture* of living R cells and heat-killed S cells contract the disease—they often die of pneumonia (**Figure 1.2**). Bacteria isolated from blood samples of these dead mice produce S cultures with a capsule typical of the injected S cells, even though the injected S cells had been killed by heat. Evidently, the injected material from the dead S cells includes a substance that can be transferred to living R cells and confer the ability to resist the immunological system of the mouse and cause pneumonia. In other words, the R bacteria can be changed—or undergo **transformation**—into S bacteria. Furthermore, the new characteristics are inherited by descendants of the transformed bacteria.

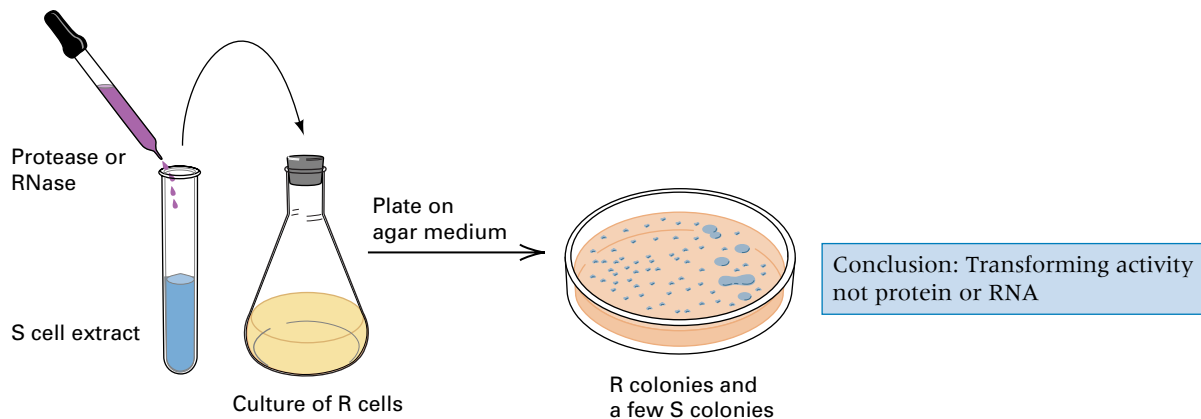
Transformation in *Streptococcus* was originally discovered in 1928, but it was not

until 1944 that the chemical substance responsible for changing the R cells into S cells was identified. In a milestone experiment, Oswald Avery, Colin MacLeod, and Maclyn McCarty showed that the substance causing the transformation of R cells into S cells was DNA. In doing these experiments, they first had to develop chemical procedures for isolating almost pure DNA from cells, which had never been done before. When they added DNA isolated from S cells to growing cultures of R cells, they observed transformation: A few cells of type S cells were produced. Although the DNA preparations contained traces of protein and RNA (ribonucleic acid, an abundant cellular macromolecule chemically related to DNA), the transforming activity was not altered by treatments that destroyed either protein or RNA. However, treatments that destroyed DNA eliminated the transforming activity (**Figure 1.3**). These experiments implied that the substance responsible for genetic transformation was the DNA of the cell—hence that DNA is the genetic material.

(A) The transforming activity in S cells is not destroyed by heat.



(B) The transforming activity is not destroyed by either protease or RNase.



(C) The transforming activity is destroyed by DNase.

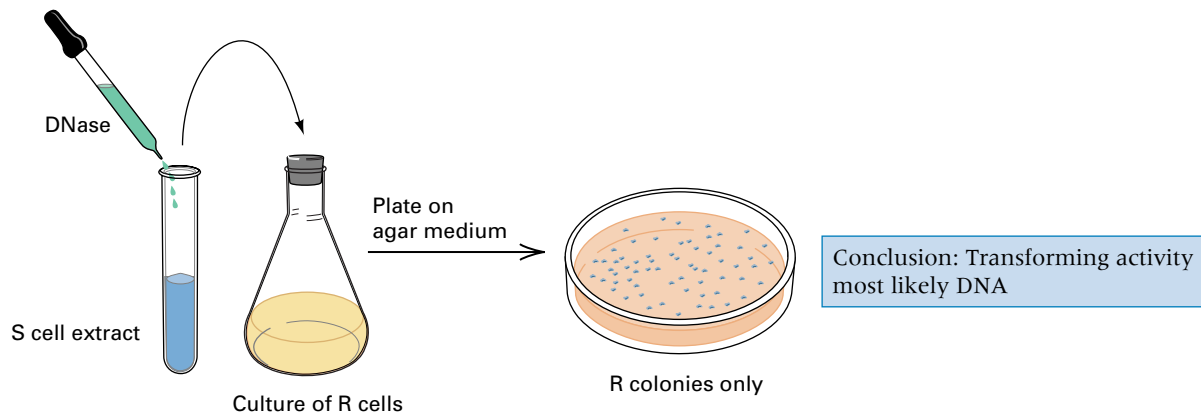


Figure 1.3 A diagram of the Avery–MacLeod–McCarty experiment that demonstrated that DNA is the active material in bacterial transformation. (A) Purified DNA extracted from heat-killed S cells can convert some living R cells into S cells, but the material may still contain undetectable traces of protein and/or RNA. (B) The transforming activity is not destroyed by either protease or RNase. (C) The transforming activity is destroyed by DNase and so probably consists of DNA.

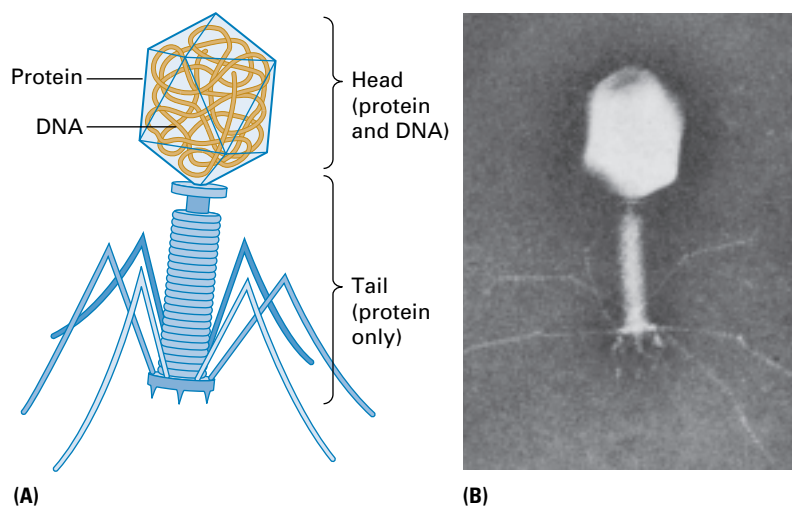


Figure 1.4 (A) Drawing of *E. coli* phage T2, showing various components. The DNA is confined to the interior of the head. (B) An electron micrograph of phage T4, a closely related phage. [Electron micrograph courtesy of Robley Williams.]

Genetic Role of DNA in Bacteriophage

A second pivotal finding was reported by Alfred Hershey and Martha Chase in 1952. They studied cells of the intestinal bacterium *Escherichia coli* after infection by the virus T2. A virus that attacks bacterial cells is called a **bacteriophage**, a term often shortened to **phage**. *Bacteriophage* means “bacteria-eater.” The structure of a bacteriophage T2 particle is illustrated in **Figure 1.4**. It is exceedingly small, yet it has a complex structure composed of head (which contains the phage DNA), collar, tail, and tail fibers. (The head of a human sperm is about 30–50 times larger in both length and width than the head of T2.) Hershey and Chase were already aware that T2 infection proceeds via the attachment of a phage particle by the tip of its tail to the bacterial cell wall, entry of phage material into the cell, multiplication of this material to form a hundred or more progeny phage, and release of the progeny phage by bursting (lysis) of the bacterial host cell. They also knew that T2 particles were composed of DNA and protein in approximately equal amounts.

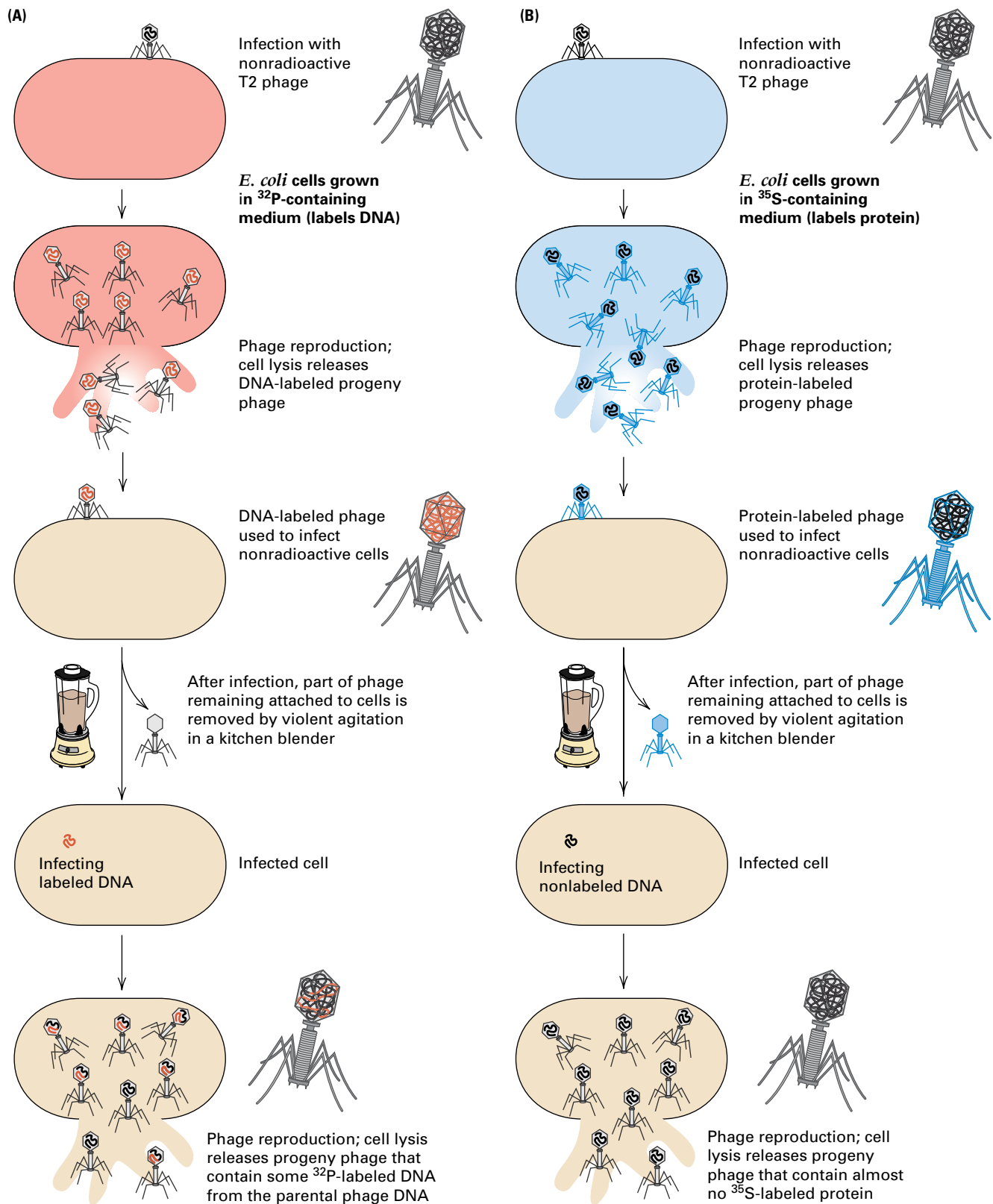
Because DNA contains phosphorus but no sulfur, whereas most proteins contain sulfur but no phosphorus, it is possible to label DNA and proteins differentially by using

radioactive isotopes of the two elements. Hershey and Chase produced particles containing radioactive DNA by infecting *E. coli* cells that had been grown for several generations in a medium that included ^{32}P (a radioactive isotope of phosphorus) and then collecting the phage progeny. Other particles containing labeled proteins were obtained in the same way, by using medium that included ^{35}S (a radioactive isotope of sulfur).

In the experiments summarized in **Figure 1.5**, nonradioactive *E. coli* cells were infected with phage labeled with either ^{32}P (part A) or ^{35}S (part B) in order to follow the DNA and proteins individually. Infected cells were separated from unattached phage particles by centrifugation, resuspended in fresh medium, and then swirled violently in a kitchen blender to shear attached phage material from the cell surfaces. This treatment was found to have no effect on the subsequent course of the infection, which implies that the phage genetic material must enter the infected cells very soon after phage attachment. The kitchen blender turned out to be the critical piece of equipment. Other methods had been tried to tear the phage heads from the bacterial cell surface, but nothing had worked reliably. Hershey later explained, “We tried various grinding arrangements, with results that weren’t very encouraging. When Margaret McDonald loaned us her kitchen blender, the experiment promptly succeeded.”

After the phage heads were removed by the blender treatment, the infected bacteria were examined. Most of the radioactivity from ^{32}P -labeled phage was found to be associated with the bacteria, whereas only a small fraction of the ^{35}S radioactivity was present in the infected cells. The retention of most of the labeled DNA, contrasted with the loss of most of the labeled protein, implied that a T2 phage transfers most of its DNA, but very little of its protein, to the cell it infects. The critical finding (**Figure 1.5**)

Figure 1.5 (on facing page) The Hershey–Chase (“blender”) experiment demonstrating that DNA, not protein, is responsible for directing the reproduction of phage T2 in infected *E. coli* cells. (A) Radioactive DNA is transmitted to progeny phage in substantial amounts. (B) Radioactive protein is transmitted to progeny phage in negligible amounts.



Conclusion: DNA from an infecting parental phage is inherited in the progeny phage

Alfred D. Hershey and Martha Chase 1952

Cold Spring Harbor Laboratories,
Cold Spring Harbor, New York
*Independent Functions of Viral Protein
and Nucleic Acid in Growth of
Bacteriophage*

Published a full eight years after the paper of Avery, MacLeod, and McCarty, the experiments of Hershey and Chase get equal billing. Why? Some historians of science suggest that the Avery et al. experiments were "ahead of their time." Others suggest that Hershey had special standing because he was a member of the "in group" of phage molecular geneticists. Max Delbrück was the acknowledged leader of this group, with Salvador Luria close behind. (Delbrück, Luria, and Hershey shared a 1969 Nobel Prize.) Another possible reason is that whereas the experiments of Avery et al. were feats of strength in biochemistry, those of Hershey and Chase were quintessentially genetic. Which macromolecule gets into the hereditary action, and which does not? Buried in the middle of this paper, and retained in the excerpt, is a sentence admitting that an earlier publication by the researchers was a misinterpretation of their preliminary results. This shows that even first-rate scientists, then and now, are sometimes misled by their preliminary data. Hershey later explained, "We tried various grinding arrangements, with results that weren't very encouraging. When

Margaret McDonald loaned us her kitchen blender the experiment promptly succeeded."

The work [of others] has shown that bacteriophages T2, T3, and T4 multiply in the bacterial cell in a non-infective [immature] form. Little else is known about the vegetative [growth] phase of these viruses. The experiments reported in this paper show that one of the first steps in the growth of T2 is the release from its protein coat of the nucleic acid of the virus particle, after which the bulk of the sulfur-containing protein has no further function. . . . Anderson has obtained electron micrographs indicating that phage T2 attaches to bacteria by its tail. . . . It ought to be a simple matter to break the empty phage coats off the infected bacteria, leaving the phage DNA inside the cells. . . . When a suspension of cells with ^{35}S - or ^{32}P -labeled phage was spun in a blender at 10,000 revolutions per minute, . . . 75 to 80 percent of the phage sulfur can be stripped from the infected cells. . . . These facts show that the bulk of the phage sulfur remains at the cell surface during infection. . . . Little or no ^{35}S is contained in the mature phage progeny. . . . Identical experiments starting with phage labeled with ^{32}P show that phosphorus is trans-

ferred from parental to progeny phage at yields of about 30 phage per infected bacterium. . . . [Incomplete separation of phage heads] explains a mistaken preliminary report of the transfer of ^{35}S from parental to progeny phage. . . . The following questions remain unanswered. (1) Does any sulfur-free phage material other than DNA enter the cell? (2) If so, is it transferred to the phage progeny? (3) Is the transfer of phosphorus to progeny direct or indirect? . . . Our experiments show clearly

Our experiments show clearly that a physical separation of the phage T2 into genetic and nongenetic parts is possible.

that a physical separation of the phage T2 into genetic and nongenetic parts is possible. The chemical identification of the genetic part must wait until some of the questions above have been answered. . . . The sulfur-containing protein of resting phage particles is confined to a protective coat that is responsible for the adsorption to bacteria, and functions as an instrument for the injection of the phage DNA into the cell. This protein probably has no function in the growth of the intracellular phage. The DNA has some function. Further chemical inferences should not be drawn from the experiments presented.

Source: *Journal of General Physiology* 36: 39–56

was that about 50 percent of the transferred ^{32}P -labeled DNA, but less than 1 percent of the transferred ^{35}S -labeled protein, was inherited by the *progeny* phage particles. Hershey and Chase interpreted this result to mean that the genetic material in T2 phage is DNA.

The experiments of Avery, MacLeod, and McCarty and those of Hershey and

Chase are regarded as classics in the demonstration that genes consist of DNA. At the present time, the equivalent of the transformation experiment is carried out daily in many research laboratories throughout the world, usually with bacteria, yeast, or animal or plant cells grown in culture. These experiments indicate that DNA is the genetic material in these organisms as well as

in phage T2. Although there are no known exceptions to the generalization that DNA is the genetic material in all cellular organisms and many viruses, in a few types of viruses the genetic material consists of RNA.

1.2 DNA Structure: The Double Helix

The inference that DNA is the genetic material still left many questions unanswered. How is the DNA in a gene duplicated when a cell divides? How does the DNA in a gene control a hereditary trait? What happens to the DNA when a mutation (a change in the DNA) takes place in a gene? In the early 1950s, a number of researchers began to try to understand the detailed molecular structure of DNA in hopes that the structure alone would suggest answers to these questions. In 1953 James Watson and Francis Crick at Cambridge University proposed the first essentially correct three-dimensional structure of the DNA molecule. The structure was dazzling in its elegance and revolutionary in suggesting how DNA duplicates itself, controls hereditary traits, and undergoes mutation. Even while their tin-and-wire model of the DNA molecule was still incomplete, Crick would visit his favorite pub and exclaim “we have discovered the secret of life.”

In the Watson–Crick structure, DNA consists of two long chains of subunits, each twisted around the other to form a double-stranded helix. The double helix is right-handed, which means that as one looks along the barrel, each chain follows a clockwise path as it progresses. You can visualize the right-handed coiling in part A of **Figure 1.6** if you imagine yourself looking up into the structure from the bottom. The dark spheres outline the “backbone” of each individual strand, and they coil in a clockwise direction. The subunits of each strand are **nucleotides**, each of which contains any one of four chemical constituents called **bases** attached to a phosphorylated molecule of the 5-carbon sugar **deoxyribose**. The four bases in DNA are

- Adenine (A)
- Guanine (G)
- Thymine (T)
- Cytosine (C)

The chemical structures of the nucleotides and bases need not concern us at this time.

They are examined in Chapter 2. A key point for our present purposes is that the bases in the double helix are paired as shown in Figure 1.6B. That is:

At any position on the paired strands of a DNA molecule, if one strand has an A, then the partner strand has a T; and if one strand has a G, then the partner strand has a C.

The pairing between A and T and between G and C is said to be **complementary**; the complement of A is T, and the complement of G is C. The complementary pairing means that each base along one strand of the DNA is matched with a base in the opposite position on the other strand. Furthermore:

Nothing restricts the sequence of bases in a single strand, so any sequence could be present along one strand.

This principle explains how only four bases in DNA can code for the huge amount of information needed to make an organism. It

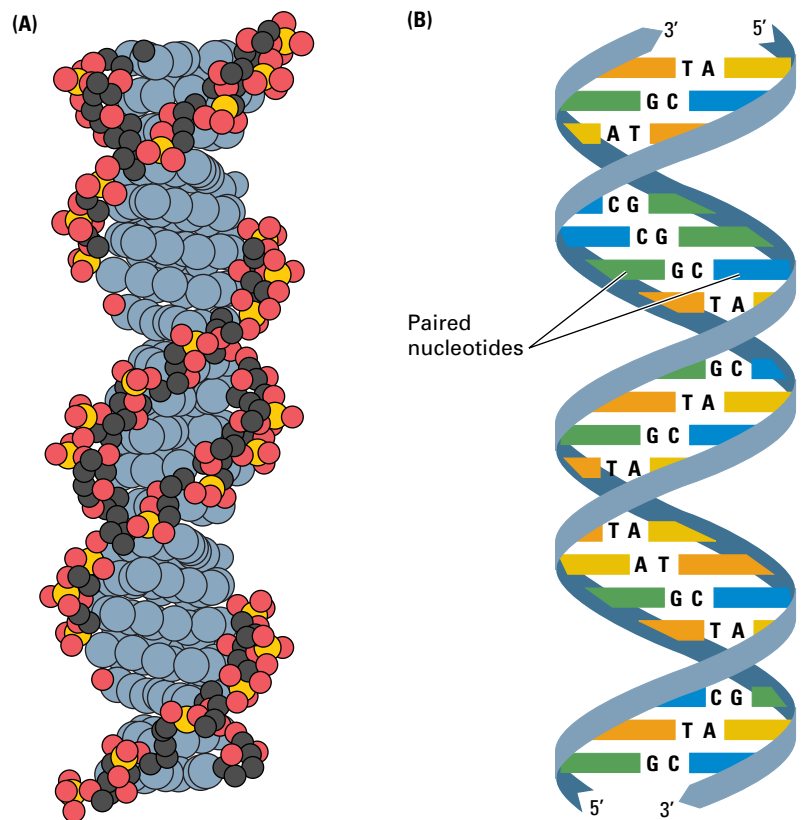


Figure 1.6 Molecular structure of the DNA double helix in the standard “B form.” (A) A space-filling model, in which each atom is depicted as a sphere. (B) A diagram highlighting the helical strands around the outside of the molecule and the A–T and G–C base pairs inside.

is the *sequence* of bases along the DNA that encodes the genetic information, and the sequence is completely unrestricted.

The complementary pairing is also called **Watson–Crick pairing**. In the three-dimensional structure in Figure 1.6A, the base pairs are represented by the lighter spheres filling the interior of the double helix. The base pairs lie almost flat, stacked on top of one another perpendicular to the long axis of the double helix, like pennies in a roll. When discussing a DNA molecule, biologists frequently refer to the individual strands as **single-stranded DNA** and to the double helix as **double-stranded DNA** or **duplex DNA**.

Each DNA strand has a **polarity**, or directionality, like a chain of circus elephants linked trunk to tail. In this analogy, each elephant corresponds to one nucleotide along the DNA strand. The polarity is determined by the direction in which the nucleotides are pointing. The “trunk” end of the strand is called the *3' end* of the strand, and the “tail” end is called the *5' end*. In double-stranded DNA, the paired strands are oriented in opposite directions, the 5' end of one strand aligned with the 3' end of the other. The molecular basis of the polarity, and the reason for the opposite orientation of the strands in duplex DNA, are explained in Chapter 2. In illustrating DNA molecules in this book, we use an arrow-like ribbon to represent the backbone, and we use tabs jutting off the ribbon to represent the nucleotides. The polarity of a DNA strand is indicated by the direction of the arrow-like ribbon. The tail of the arrow represents the 5' end of the DNA strand, the head the 3' end.

Beyond the most optimistic hopes, knowledge of the structure of DNA immediately gave clues to its function:

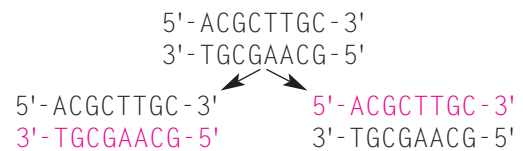
1. The sequence of bases in DNA could be copied by using each of the separate “partner” strands as a pattern for the creation of a new partner strand with a complementary sequence of bases.
2. The DNA could contain genetic information in coded form in the sequence of bases, analogous to letters printed on a strip of paper.
3. Changes in genetic information (mutations) could result from errors in copying in which the base sequence of the DNA became altered.

In the remainder of this chapter, we discuss some of the implications of these clues.

1.3 An Overview of DNA Replication

Watson and Crick noted that the structure of DNA itself suggested a mechanism for its replication. “It has not escaped our notice,” they wrote, “that the specific base pairing we have postulated immediately suggests a copying mechanism.” The copying process in which a single DNA molecule becomes two identical molecules is called **replication**. The replication mechanism that Watson and Crick had in mind is illustrated in **Figure 1.7**.

As shown in part A of Figure 1.7, the strands of the original (parent) duplex separate, and each individual strand serves as a pattern, or **template**, for the synthesis of a new strand (replica). The replica strands are synthesized by the addition of successive nucleotides in such a way that each base in the replica is complementary (in the Watson–Crick pairing sense) to the base across the way in the template strand (Figure 1.7B). Although the mechanism in Figure 1.7 is simple in principle, it is a complex process that is fraught with geometrical problems and requires a variety of enzymes and other proteins. The details are examined in Chapter 6. The end result of replication is that a single double-stranded molecule becomes replicated into two copies with identical sequences:



Here the bases in the newly synthesized strands are shown in red. In the duplex on the left, the top strand is the template from the parental molecule and the bottom strand is newly synthesized; in the duplex on the right, the bottom strand is the template from the parental molecule and the top strand is newly synthesized. Note in Figure 1.7B that in the synthesis of each new strand, new nucleotides are added only to the 3' end of the growing chain:

The obligatory elongation of a DNA strand only at the 3' end is an essential feature of DNA replication.

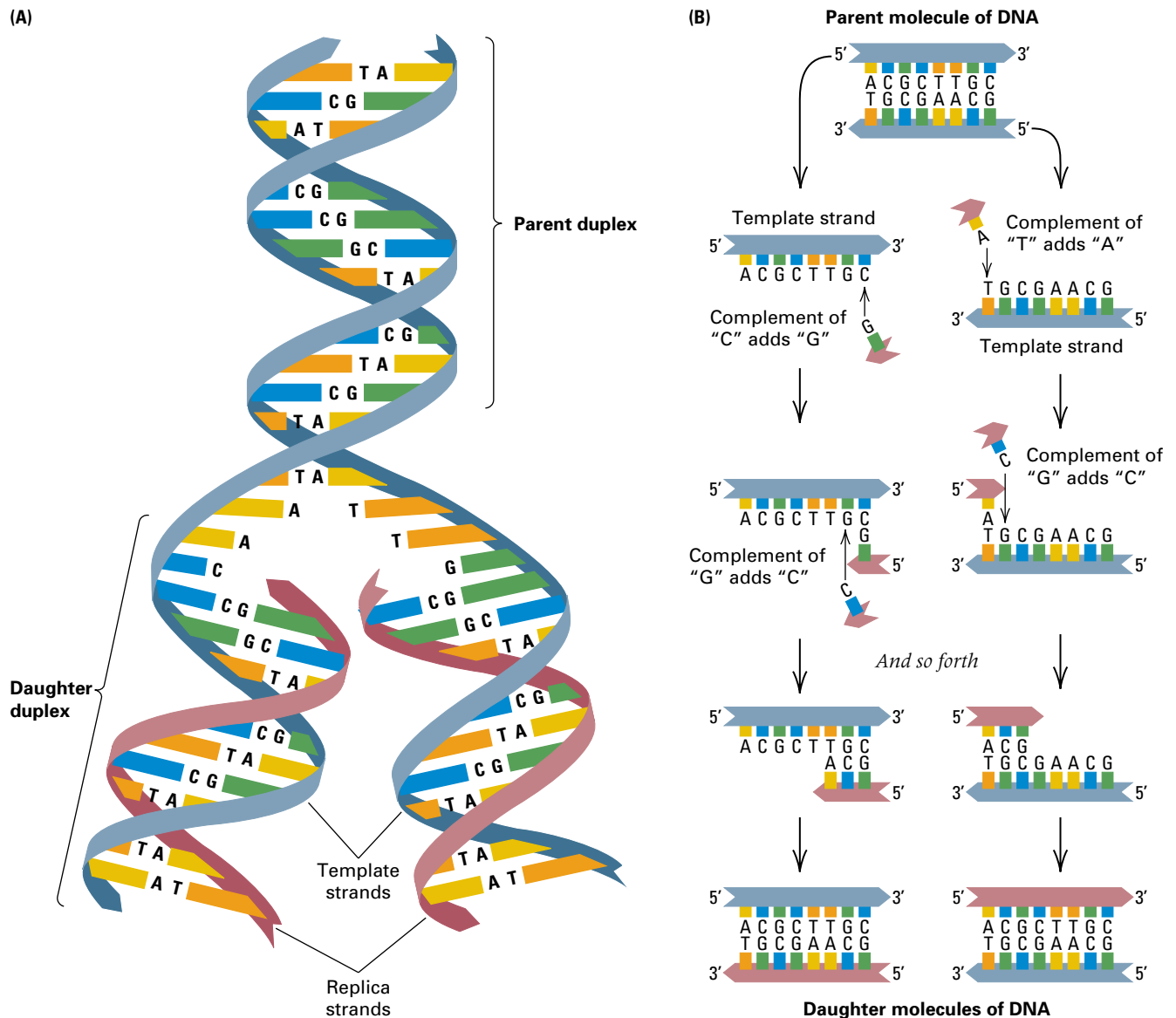


Figure 1.7 Replication of DNA. (A) Replication of a DNA duplex as originally envisioned by Watson and Crick. As the parental strands separate, each parental strand serves as a template for the formation of a new daughter strand by means of A–T and G–C base pairing. (B) Greater detail showing how each of the parental strands serves as a template for the production of a complementary daughter strand, which grows in length by the successive addition of single nucleotides to the 3' end.

1.4 Genes and Proteins

Now that we have some basic understanding of the structural makeup of the genetic blueprint, how does this developmental plan become a complex living organism? If the code is thought of as a string of letters on a sheet of paper, then the genes are made up of distinct words that form sentences and paragraphs that give meaning to the pattern of letters. What is created from

the complex and diverse DNA codes is protein, a class of macromolecules that carries out most of the activities in the cell. Cells are largely made up of proteins: structural proteins that give the cell rigidity and mobility, proteins that form pores in the cell membrane to control the traffic of small molecules into and out of the cell, and receptor proteins that regulate cellular activities in response to molecular signals from the growth medium or from other cells.

Proteins are also responsible for most of the metabolic activities of cells. They are essential for the synthesis and breakdown of organic molecules and for generating the chemical energy needed for cellular activities. In 1878 the term **enzyme** was introduced to refer to the biological catalysts that accelerate biochemical reactions in cells. By 1900, thanks largely to the work of the German biochemist Emil Fischer, enzymes were shown to be proteins. As often happens in science, nature's "mistakes" provide clues as to how things work. Such was the case in establishing a relationship between genes and disease, because a "mistake" in a gene (a mutation) can result in a "mistake" (lack of function) in the corresponding protein. This provided a fruitful avenue of research for the study of genetics.

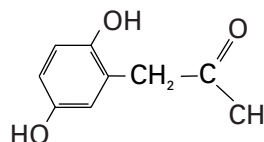
Inborn Errors of Metabolism as a Cause of Hereditary Disease

It was at the turn of the twentieth century that the British physician Archibald Garrod realized that certain heritable diseases followed the rules of transmission that Mendel had described for his garden peas. In 1908 Garrod gave a series of lectures in which he proposed a fundamental hypothesis about the relationship between heredity, enzymes, and disease:

Any hereditary disease in which cellular metabolism is abnormal results from an inherited defect in an enzyme.

Such diseases became known as **inborn errors of metabolism**, a term still in use today.

Garrod studied a number of inborn errors of metabolism in which the patients excreted abnormal substances in the urine. One of these was **alkaptonuria**. In this case, the abnormal substance excreted is **homogentisic acid**:



An early name for homogentisic acid was *alkapton*, hence the name *alkaptonuria* for the disease. Even though alkaptonuria is

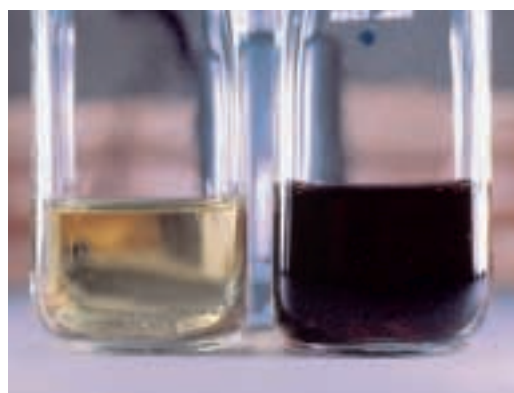


Figure 1.8 Urine from a person with alkaptonuria turns black because of the oxidation of the homogentisic acid that it contains. [Courtesy of Daniel De Aguiar.]

rare, with an incidence of about one in 200,000 people, it was well known even before Garrod studied it. The disease itself is relatively mild, but it has one striking symptom: The urine of the patient turns black because of the oxidation of homogentisic acid (**Figure 1.8**). This is why alkaptonuria is also called *black urine disease*. An early case was described in the year 1649:

The patient was a boy who passed black urine and who, at the age of fourteen years, was submitted to a drastic course of treatment that had for its aim the subduing of the fiery heat of his viscera, which was supposed to bring about the condition in question by charring and blackening his bile. Among the measures prescribed were bleedings, purgation, baths, a cold and watery diet, and drugs galore. None of these had any obvious effect, and eventually the patient, who tired of the futile and superfluous therapy, resolved to let things take their natural course. None of the predicted evils ensued. He married, begat a large family, and lived a long and healthy life, always passing urine black as ink. (Recounted by Garrod, 1908.)

Garrod was primarily interested in the biochemistry of alkaptonuria, but he took note of family studies that indicated that the disease was inherited as though it were due to a defect in a single gene. As to the biochemistry, he deduced that the problem in alkaptonuria was the patients' inability to break down the phenyl ring of six carbons that is present in homogentisic acid. Where does this ring come from? Most animals

obtain it from foods in their diet. Garrod proposed that homogentisic acid originates as a breakdown product of two amino acids, phenylalanine and tyrosine, which also contain a phenyl ring. An **amino acid** is one of the “building blocks” from which proteins are made. Phenylalanine and tyrosine are constituents of normal proteins. The scheme that illustrates the relationship between the molecules is shown in **Figure 1.9**. Any such sequence of biochemical reactions is called a **biochemical pathway** or a **metabolic pathway**. Each arrow in the pathway represents a single step depicting the transition from the “input” or **substrate molecule**, shown at the head of the arrow, to the “output” or **product molecule**, shown at the tip. Biochemical pathways are usually oriented either vertically with the arrows pointing down, as in Figure 1.9, or horizontally, with the arrows pointing from left to right. Garrod did not know all of the details of the pathway in Figure 1.9, but he did understand that the key step in the breakdown of homogentisic acid is the breaking open of the phenyl ring and that the phenyl ring in homogentisic acid comes from dietary phenylalanine and tyrosine.

What allows each step in a biochemical pathway to occur? Garrod correctly surmised that each step requires a specific enzyme to catalyze the reaction for the chemical transformation. Persons with an inborn error of metabolism, such as alkaptonuria, have a defect in a single step of a metabolic pathway because they lack a functional enzyme for that step. When an enzyme in a pathway is defective, the pathway is said to have a **block** at that step. One frequent result of a blocked pathway is that the substrate of the defective enzyme accumulates. Observing the accumulation of homogentisic acid in patients with alkaptonuria, Garrod proposed that there must be an enzyme whose function is to open the phenyl ring of homogentisic acid and that this enzyme is missing in these patients. Isolation of the enzyme that opens the phenyl ring of homogentisic acid was not actually achieved until 50 years after Garrod’s lectures. In normal people it is found in cells of the liver, and just as Garrod had predicted, the enzyme is defective in patients with alkaptonuria.

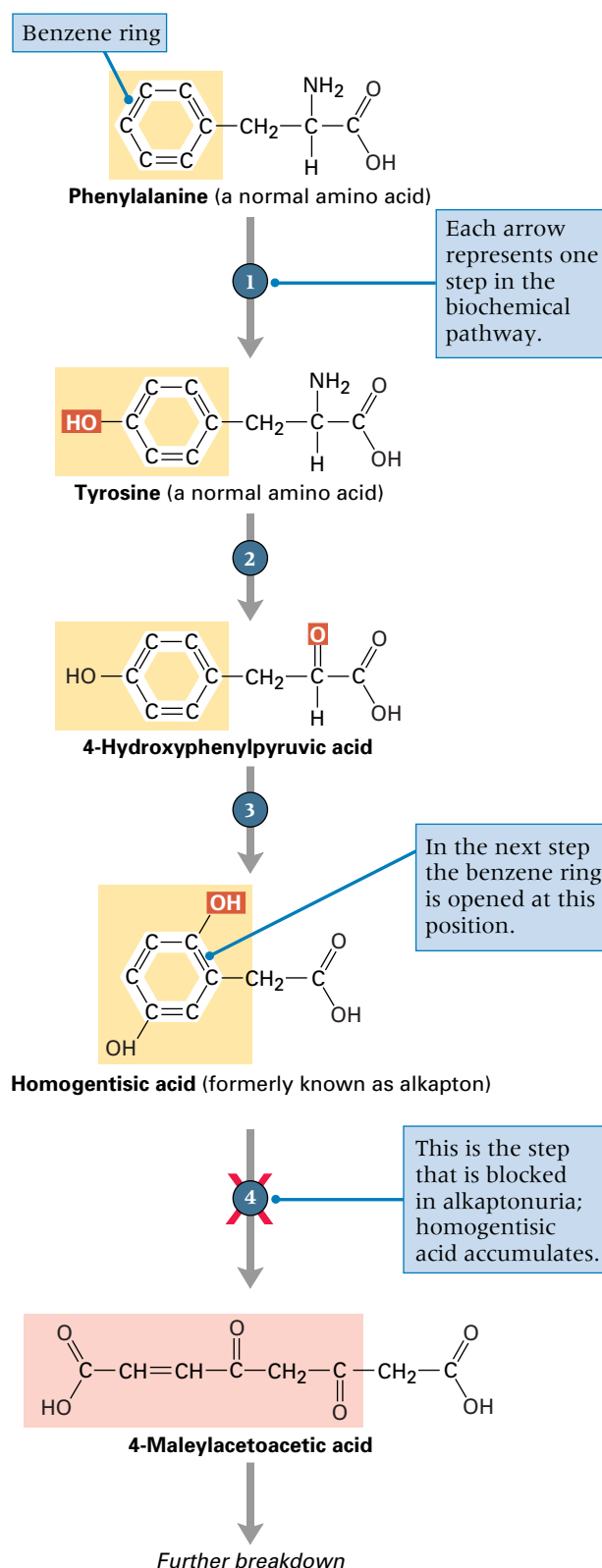


Figure 1.9 Metabolic pathway for the breakdown of phenylalanine and tyrosine. Each step in the pathway, represented by an arrow, requires a specific enzyme to catalyze the reaction. The key step in the breakdown of homogentisic acid is the breaking open of the phenyl ring.

Archibald E. Garrod 1908

St. Bartholomew's Hospital,
London, England

Inborn Errors of Metabolism

Although he was a distinguished physician, Garrod's lectures on the relationship between heredity and congenital defects in metabolism had no impact when they were delivered. The important concept that one gene corresponds to one enzyme (the "one gene—one enzyme hypothesis") was developed independently in the 1940s by George W. Beadle and Edward L. Tatum, who used the bread mold *Neurospora crassa* as their experimental organism. When Beadle finally became aware of Inborn Errors of Metabolism, he was generous in praising it. This excerpt shows Garrod at his best, interweaving history, clinical medicine, heredity, and biochemistry in his account of alkaptonuria. The excerpt also illustrates how the severity of a genetic disease depends on its social context. Garrod writes as though alkaptonuria were a harmless curiosity. This is indeed largely true when the life expectancy is short. With today's longer life span, alkaptonuria patients accumulate the dark pigment in their cartilage and joints and eventually develop severe arthritis.

To students of heredity the inborn errors of metabolism offer a promising field of investigation. . . . It was pointed out [by others] that the mode of incidence of alkaptonuria finds a ready explanation if the anomaly be regarded as a rare recessive character in the Mendelian sense. . . . Of the cases of alkaptonuria a very large proportion have been in the children of first cousin marriages. . . . It is also noteworthy that, if one takes families with five or more children [with both parents normal and at least one child affected with alkaptonuria], the totals work out in strict conformity to Mendel's law, i.e. 57 [normal children]:19 [affected children] in the proportions 3:1. . . . Of inborn errors of metabolism, alkaptonuria is that of which we know most. In itself it is a trifling matter, inconvenient rather than harmful. . . . Indications of the anomaly may be detected in early medical writings, such as that in 1584 of a schoolboy who, although he enjoyed good health, continuously excreted black urine; and that in 1609 of a monk who exhibited a similar peculiarity and stated that he had done so all his life. . . . There are no sufficient grounds [for doubting that the blackening substance

in the urine originally called alkapton] is homogentisic acid, the excretion of which is the essential feature of the alkaptonuric. . . . Homogentisic acid is a product of normal metabolism. . . . The

most likely sources of the benzene ring in homogentisic acid are phenylalanine and tyrosine, [because when these amino acids are administered to an alkaptonuric] they cause a very conspicuous increase in the output of homogentisic acid. . . . Where the alkaptonuric differs from the normal individual is in having no power of destroying homogentisic acid when formed—in

other words of breaking up the benzene ring of that compound. . . . We may further conceive that the splitting of the benzene ring in normal metabolism is the work of a special enzyme and that in congenital alkaptonuria this enzyme is wanting.

We may further conceive that the splitting of the benzene ring in normal metabolism is the work of a special enzyme and that in congenital alkaptonuria this enzyme is wanting.

Source: Originally published in London, England, by the Oxford University Press. Excerpts from the reprinted edition in Harry Harris. 1963. *Garrod's Inborn Errors of Metabolism*. London, England: Oxford University Press.

The pathway for the breakdown of phenylalanine and tyrosine, as it is understood today, is shown in **Figure 1.10**. In this figure the emphasis is on the enzymes rather than on the structures of the **metabolites**, or small molecules, on which the enzymes act. Each step in the pathway requires the presence of a particular enzyme that catalyzes that step. Although Garrod knew only about alkaptonuria, in which the defective enzyme is homogentisic acid 1,2 dioxygenase, we now know the

clinical consequences of defects in the other enzymes. Unlike alkaptonuria, which is a relatively benign inherited disease, the others are very serious. The condition known as **phenylketonuria (PKU)** results from the absence of (or a defect in) the enzyme **phenylalanine hydroxylase (PAH)**. When this step in the pathway is blocked, phenylalanine accumulates. The excess phenylalanine is broken down into harmful metabolites that cause defects in myelin formation that damage a child's developing

nervous system and lead to severe mental retardation.

However, if PKU is diagnosed in children soon enough after birth, they can be placed on a specially formulated diet low in phenylalanine. The child is allowed only as much phenylalanine as can be used in the synthesis of proteins, so excess phenylalanine does not accumulate. The special diet is very strict. It excludes meat, poultry, fish, eggs, milk and milk products, legumes, nuts, and bakery goods manufactured with regular flour. These foods are replaced by an expensive synthetic formula. With the special diet, however, the detrimental effects of excess phenylalanine on mental development can largely be avoided, although in adult women with PKU who are pregnant, the fetus is at risk. In many countries, including the United States, all newborn babies have their blood tested for chemical signs of PKU. Routine screening is cost-effective because PKU is relatively common. In the United States, the incidence is about 1 in 8000 among Caucasian births. The disease is less common in other ethnic groups.

In the metabolic pathway in Figure 1.10, defects in the breakdown of tyrosine or of 4-hydroxyphenylpyruvic acid lead to types of tyrosinemia. These are also severe diseases. Type II is associated with skin lesions and mental retardation, Type III with severe liver dysfunction.

Mutant Genes and Defective Proteins

It follows from Garrod's work that a defective enzyme results from a mutant gene, but how? Garrod did not speculate. For all he knew, genes *were* enzymes. This would have been a logical hypothesis at the time. We now know that the relationship between genes and enzymes is somewhat indirect. With a few exceptions, each enzyme is *encoded* in a particular sequence of nucleotides present in a region of DNA. The DNA region that codes for the enzyme, as well as adjacent regions that regulate when and in which cells the enzyme is produced, make up the "gene" that encodes the enzyme.

The genes for the enzymes in the biochemical pathway in Figure 1.10 have all been identified and the nucleotide sequence of the DNA determined. In the following list, and throughout this book, we

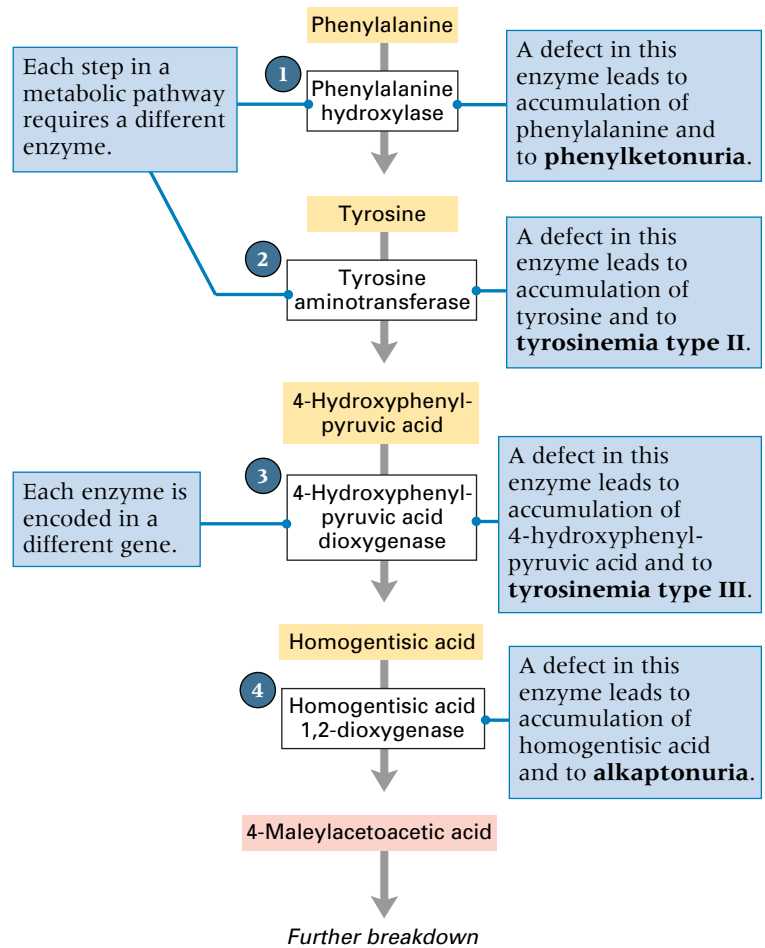


Figure 1.10 Inborn errors of metabolism that affect the breakdown of phenylalanine and tyrosine. An inherited disease results when any of the enzymes is missing or defective. Alkaptonuria results from a mutant homogentisic acid 1,2 dioxygenase phenylketonuria results from a mutant phenylalanine hydroxylase.

use the standard typographical convention that *genes* are written in *italic* type, whereas gene products are not printed in italics. This convention is convenient, because it means that the protein product of a gene can be represented with the same symbol as the gene itself, but whereas the gene symbol is in italics, the protein symbol is not.

- The gene *PAH* on the long arm of chromosome 12 encodes phenylalanine hydroxylase (PAH).
- The gene *TAT* on the long arm of chromosome 16 encodes tyrosine aminotransferase (TAT).
- The gene *HPD* on the long arm of chromosome 12 encodes 4-hydroxyphenylpyruvic acid dioxygenase (HPD).

- The gene *HGD* on the long arm of chromosome 3 encodes homogentisic acid 1,2 dioxygenase (HGD).

Next we turn to the issue of *how* genes code for enzymes and other proteins.

1.5 Gene Expression: The Central Dogma

Watson and Crick were correct in proposing that the genetic information in DNA is contained in the sequence of bases in a manner analogous to letters printed on a strip of paper. In a region of DNA that directs the synthesis of a protein, the genetic code for the protein is contained in only one strand, and it is decoded in a linear order. A typical protein is made up of one or more polypeptide

chains; each **polypeptide chain** consists of a linear sequence of amino acids connected end to end. For example, the enzyme PAH consists of four identical polypeptide chains, each 452 amino acids in length. In the decoding of DNA, each successive “code word” in the DNA specifies the next amino acid to be added to the polypeptide chain as it is being made. The amount of DNA required to code for the polypeptide chain of PAH is therefore $452 \times 3 = 1356$ nucleotide pairs. The entire gene is very much longer—about 90,000 nucleotide pairs. Only 1.5 percent of the gene is devoted to coding for the amino acids. The noncoding part includes some sequences that control the activity of the gene, but it is not known how much of the gene is involved in regulation.

There are 20 different amino acids. Only four bases code for these 20 amino acids, with each “word” in the genetic code consisting of three adjacent bases. For example, the base sequence ATG specifies the amino acid methionine (Met), TCC specifies serine (Ser), ACT specifies threonine (Thr), and GCG specifies alanine (Ala). There are 64 possible three-base combinations but only 20 amino acids because some combinations code for the same amino acid. For example, TCT, TCC, TCA, TCG, AGT, and AGC all code for serine (Ser), and CTT, CTC, CTA, CTG, TTA, and TTG all code for leucine (Leu). An example of the relationship between the base sequence in a DNA duplex and the amino acid sequence of the corresponding protein is shown in **Figure 1.11**. This particular DNA duplex is the human sequence that codes for the first seven amino acids in the polypeptide chain of PAH.

The scheme outlined in Figure 1.11 indicates that DNA codes for protein not directly but indirectly through the processes of *transcription* and *translation*. The indirect route of information transfer,

DNA → RNA → Protein

is known as the **central dogma** of molecular genetics. The term *dogma* means “set of beliefs”; it dates from the time the idea was put forward first as a theory. Since then the “dogma” has been confirmed experimentally, but the term persists. The central dogma is shown in **Figure 1.12**. The main concept in the central dogma is that DNA

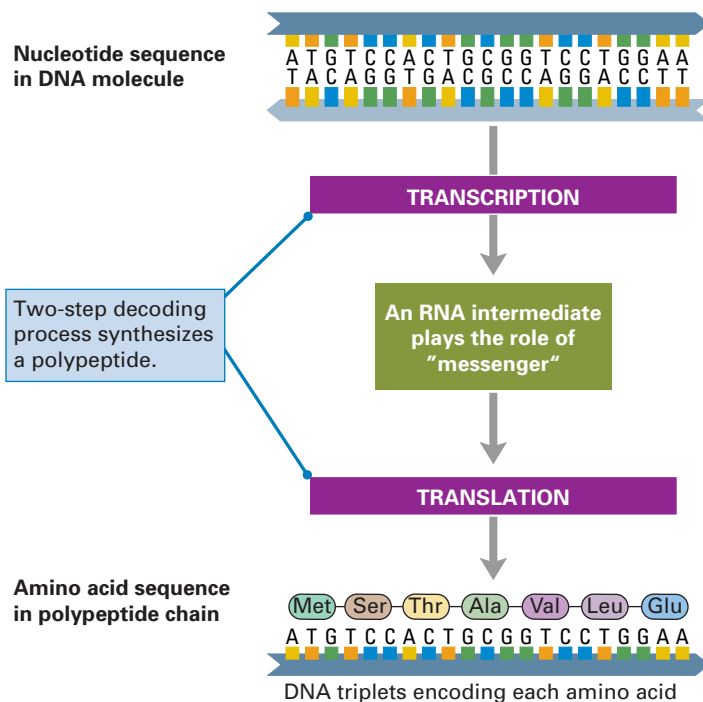


Figure 1.11 DNA sequence coding for the first seven amino acids in a polypeptide chain. The DNA sequence specifies the amino acid sequence through a molecule of RNA that serves as an intermediary “messenger.” Although the decoding process is indirect, the net result is that each amino acid in the polypeptide chain is specified by a group of three adjacent bases in the DNA. In this example, the polypeptide chain is that of phenylalanine hydroxylase (PAH).

does not code for protein directly but rather acts through an intermediary molecule of **ribonucleic acid (RNA)**. The structure of RNA is similar to, but not identical with, that of DNA. There is a difference in the sugar (RNA contains the sugar **ribose** instead of deoxyribose), RNA is usually single-stranded (not a duplex), and RNA contains the base **uracil (U)** instead of thymine (T), which is present in DNA. Actually, three types of RNA take part in the synthesis of proteins:

- A molecule of **messenger RNA (mRNA)**, which carries the genetic information from DNA and is used as a template for polypeptide synthesis. In most mRNA molecules, there is a high proportion of nucleotides that actually code for amino acids. For example, the mRNA for PAH is 2400 nucleotides in length and codes for a polypeptide of 452 amino acids; in this case, more than 50 percent of the length of the mRNA codes for amino acids.
- Several types of **ribosomal RNA (rRNA)**, which are major constituents of the cellular particles called **ribosomes** on which polypeptide synthesis takes place.
- A set of **transfer RNA (tRNA)** molecules, each of which carries a particular amino acid as well as a three-base recognition region that base-pairs with a group of three adjacent bases in the mRNA. As each tRNA participates in translation, its amino acid becomes the terminal subunit added to the length of the growing polypeptide chain. The tRNA that carries methionine is denoted tRNA^{Met} , that which carries serine is denoted tRNA^{Ser} , and so forth.

The central dogma is the fundamental principle of molecular genetics because it summarizes how the genetic information in DNA becomes expressed in the amino acid sequence in a polypeptide chain:

The sequence of nucleotides in a gene specifies the sequence of nucleotides in a molecule of messenger RNA; in turn, the sequence of nucleotides in the messenger RNA specifies the sequence of amino acids in the polypeptide chain.

Given a process as conceptually simple as DNA coding for protein, what might account for the additional complexity of RNA intermediaries? One possible reason is that an RNA intermediate gives another level for control, for example, by degrading the

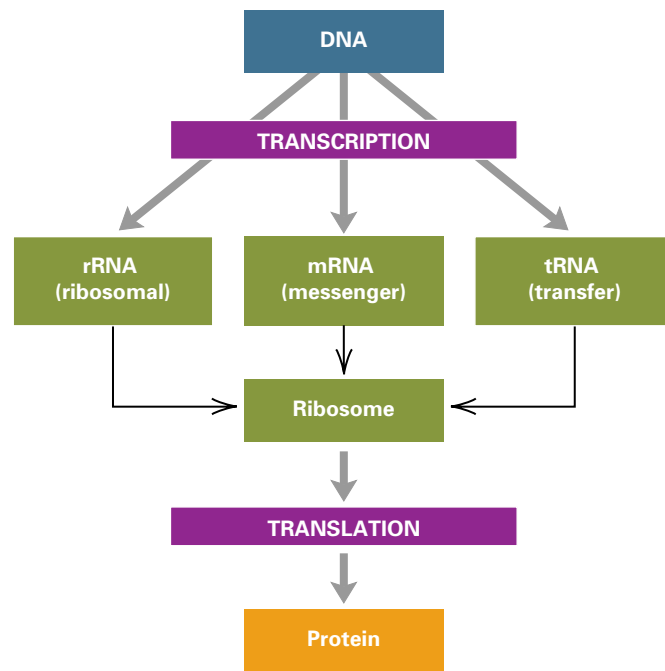


Figure 1.12 The “central dogma” of molecular genetics: DNA codes for RNA, and RNA codes for protein. The DNA → RNA step is transcription, and the RNA → protein step is translation.

mRNA for an unneeded protein. Another possible reason may be historical. RNA structure is unique in having both an informational content present in its sequence of bases and a complex, folded three-dimensional structure that endows some RNA molecules with catalytic activities. Many scientists believe that in the earliest forms of life, RNA served both for genetic information and catalysis. As evolution proceeded, the informational role was transferred to DNA and the catalytic role to protein. However, RNA became locked into its central location as a go-between in the processes of information transfer and protein synthesis. This hypothesis implies that the participation of RNA in protein synthesis is a relic of the earliest stages of evolution—a “molecular fossil.” The hypothesis is supported by a variety of observations. For example, (1) DNA replication requires an RNA molecule in order to get started (Chapter 6), (2) an RNA molecule is essential in the synthesis of the tips of the chromosomes (Chapter 8), and (3) some RNA molecules act to catalyze key reactions in protein synthesis (Chapter 11).

Transcription

The manner in which genetic information is transferred from DNA to RNA is shown in **Figure 1.13**. The DNA opens up, and one of the strands is used as a template for the synthesis of a complementary strand of RNA. (How the template strand is chosen is discussed in Chapter 11.) The process of making an RNA strand from a DNA template is **transcription**, and the RNA molecule that is made is the **transcript**. The base sequence in the RNA is complementary (in

the Watson–Crick pairing sense) to that in the DNA template, except that U (which pairs with A) is present in the RNA in place of T. The rules of base pairing between DNA and RNA are summarized in **Figure 1.14**. Each RNA strand has a polarity—a 5' end and a 3' end—and, as in the synthesis of DNA, nucleotides are added only to the 3' end of a growing RNA strand. Hence the 5' end of the RNA transcript is synthesized first, and transcription proceeds along the template DNA strand in the 3'-to-5' direction. Each gene includes nucleotide sequences that initiate and terminate transcription. The RNA transcript made from any gene begins at the initiation site in the template strand, which is located “upstream” from the amino acid-coding region, and ends at the termination site, which is located “downstream” from the amino acid-coding region. For any gene, the length of the RNA transcript is very much smaller than the length of the DNA in the chromosome. For example, the transcript of the *PAH* gene for phenylalanine hydroxylase is about 90,000 nucleotides in length, but the DNA in chromosome 12 is about 130,000,000 nucleotide pairs. In this case, the length of the *PAH* transcript is less than 0.1 percent of the length of the DNA in the chromosome. A different gene in chromosome 12 would be transcribed from a different region of the DNA molecule in chromosome 12, and perhaps from the opposite strand, but the transcribed region would again be small in comparison with the total length of the DNA in the chromosome.

Translation

The synthesis of a polypeptide under the direction of an mRNA molecule is known as **translation**. Although the sequence of bases in the mRNA codes for the sequence of amino acids in a polypeptide, the molecules that actually do the “translating” are the tRNA molecules. The mRNA molecule is translated in nonoverlapping groups of three bases called **codons**. For each codon in the mRNA that specifies an amino acid, there is one tRNA molecule containing a complementary group of three adjacent bases that can pair with those in that codon. The correct amino acid is attached to the other end of the tRNA, and when the tRNA comes into line, the amino acid to which it

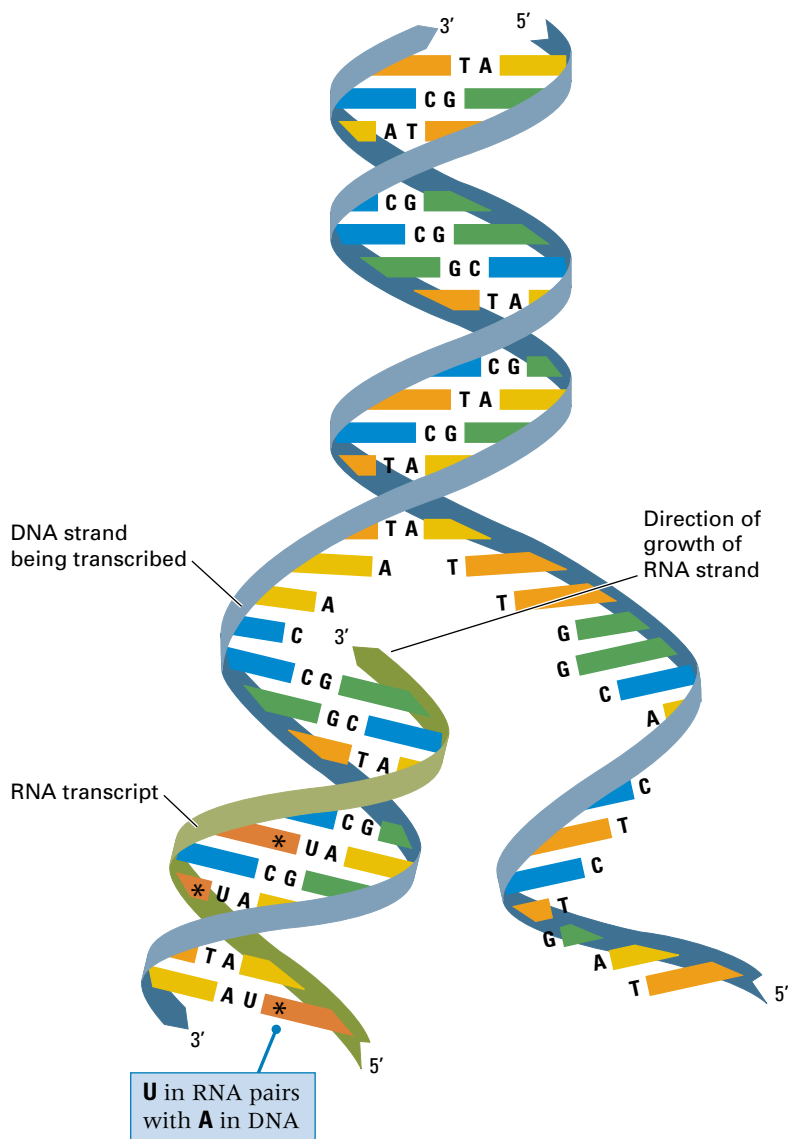


Figure 1.13 Transcription is the production of an RNA strand that is complementary in base sequence to a DNA strand. In this example, the DNA strand at the bottom is being transcribed into a strand of RNA. Note that in an RNA molecule, the base U (uracil) plays the role of T (thymine) in that it pairs with A (adenine). Each A–U pair is marked.

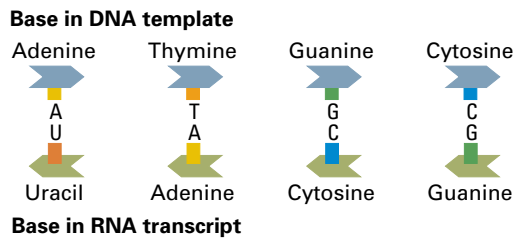


Figure 1.14 Pairing between bases in DNA and in RNA. The DNA bases A, T, G, and C pair with the RNA bases U, A, C, and G, respectively.

is attached becomes the most recent addition to the growing end of the polypeptide chain.

The role of tRNA in translation is illustrated in **Figure 1.15** and can be described as follows:

The mRNA is read codon by codon. Each codon that specifies an amino acid matches with a complementary group of three adjacent bases in a single tRNA molecule. One end of the tRNA is attached to the correct amino acid, so the correct amino acid is brought into line.

The tRNA molecules used in translation do not line up along the mRNA simultaneously as shown in **Figure 1.15**. The process of translation takes place on a ribosome,

which combines with a single mRNA and moves along it from one end to the other in steps, three nucleotides at a time (codon by codon). As each new codon comes into place, the next tRNA binds with the ribosome. Then the growing end of the polypeptide chain becomes attached to the amino acid on the tRNA. In this way, each tRNA in turn serves temporarily to hold the polypeptide chain as it is being synthesized. As the polypeptide chain is transferred from each tRNA to the next in line, the tRNA that previously held the polypeptide is released from the ribosome. The polypeptide chain elongates one amino acid at each step until any one of three particular codons specifying “stop” is encountered. At this point, synthesis of the chain of amino acids is finished, and the polypeptide chain is released from the ribosome. (This brief description of translation glosses over many of the details that are presented in Chapter 11.)

The Genetic Code

Figure 1.15 indicates that the mRNA codon AUG specifies methionine (Met) in the polypeptide chain, UCC specifies Ser (serine), ACU specifies Thr (threonine), and so on. The complete decoding table is

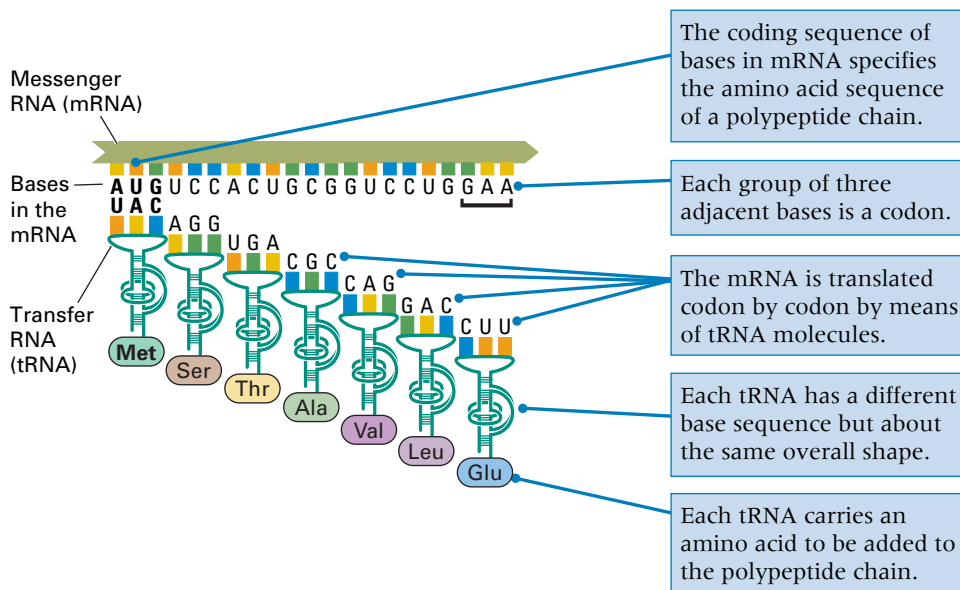


Figure 1.15 The role of messenger RNA in translation is to carry the information contained in a sequence of DNA bases to a ribosome, where it is translated into a polypeptide chain. Translation is mediated by transfer RNA (tRNA) molecules, each of which can base-pair with a group of three adjacent bases in the mRNA. Each tRNA also carries an amino acid. As each tRNA, in turn, is brought to the ribosome, the growing polypeptide chain is elongated.

Table 1.1 The standard genetic code

Second nucleotide in codon																	
U				C				A				G					
First nucleotide in codon (5' end)	U	UUU	Phe	F	Phenylalanine	UCU	Ser	S	Serine	UAU	Tyr	Y	Tyrosine	UGU	Cys	C	Cysteine
		UUC	Phe	F	Phenylalanine	UCC	Ser	S	Serine	UAC	Tyr	Y	Tyrosine	UGC	Cys	C	Cysteine
		UUA	Leu	L	Leucine	UCA	Ser	S	Serine	UAA	Termination		UGA	Termination			
		UUG	Leu	L	Leucine	UCG	Ser	S	Serine	UAG	Termination		UGG	Trp	W	Tryptophan	
	C	CUU	Leu	L	Leucine	CCU	Pro	P	Proline	CAU	His	H	Histidine	CGU	Arg	R	Arginine
		CUC	Leu	L	Leucine	CCC	Pro	P	Proline	CAC	His	H	Histidine	CGC	Arg	R	Arginine
		CUA	Leu	L	Leucine	CCA	Pro	P	Proline	CAA	Gln	Q	Glutamine	CGA	Arg	R	Arginine
		CUG	Leu	L	Leucine	CCG	Pro	P	Proline	CAG	Gln	Q	Glutamine	CGG	Arg	R	Arginine
	A	AUU	Ile	I	Isoleucine	ACU	Thr	T	Threonine	AAU	Asn	N	Asparagine	AGU	Ser	S	Serine
		AUC	Ile	I	Isoleucine	ACC	Thr	T	Threonine	AAC	Asn	N	Asparagine	AGC	Ser	S	Serine
		AUA	Ile	I	Isoleucine	ACA	Thr	T	Threonine	AAA	Lys	K	Lysine	AGA	Arg	R	Arginine
		AUG	Met	M	Methionine	ACG	Thr	T	Threonine	AAG	Lys	K	Lysine	AGG	Arg	R	Arginine
	G	GUU	Val	V	Valine	GCU	Ala	A	Alanine	GAU	Asp	D	Aspartic acid	GGU	Gly	G	Glycine
		GUC	Val	V	Valine	GCC	Ala	A	Alanine	GAC	Asp	D	Aspartic acid	GGC	Gly	G	Glycine
		GUA	Val	V	Valine	GCA	Ala	A	Alanine	GAA	Glu	E	Glutamic acid	GGA	Gly	G	Glycine
		GUG	Val	V	Valine	GCG	Ala	A	Alanine	GAG	Glu	E	Glutamic acid	GGG	Gly	G	Glycine

Codon

Three-letter and single-letter abbreviations

Codon

Three-letter and single-letter abbreviations

called the **genetic code**, and it is shown in [Table 1.1](#). For any codon, the column on the left corresponds to the first nucleotide in the codon (reading from the 5' end), the row across the top corresponds to the second nucleotide, and the column on the right corresponds to the third nucleotide. The complete codon is given in the body of the table, along with the amino acid (or translational “stop”) that the codon specifies. Each amino acid is designated by its full name and by a three-letter abbreviation as well as a single-letter abbreviation. Both types of abbreviations are used in molecular genetics. The code in Table 1.1 is the “standard” genetic code used in translation in the cells of nearly all organisms. In Chapter 11 we examine general features of the standard genetic code and the minor differences found in the genetic codes of certain organisms and cellular organelles. At this point, we are interested mainly in understanding how the genetic code is used to translate the codons in mRNA into the amino acids in a polypeptide chain.

In addition to the 61 codons that code only for amino acids, there are four codons that have specialized functions:

- The codon AUG, which specifies Met (methionine), is also the “start” codon for polypeptide synthesis. The positioning of a tRNA^{Met} bound to AUG is one of the first steps in the initiation of polypeptide synthesis, so all polypeptide chains begin with Met. (Many polypeptides have the initial Met cleaved off after translation is complete.) In most organisms, the tRNA^{Met} used for initiation of translation is the same tRNA^{Met} used to specify methionine at internal positions in a polypeptide chain.
- The codons UAA, UAG, and UGA are each a “stop” that specifies the termination of translation and results in release of the completed polypeptide chain from the ribosome. These codons do not have tRNA molecules that recognize them but are instead recognized by protein factors that terminate translation.

How the genetic code table is used to infer the amino acid sequence of a polypep-

tide chain can be illustrated by using PAH again, in particular the DNA sequence coding for amino acids 1 through 7. The DNA sequence is

5'-ATGTCCACTGCGGTCCTGGAA-3'
3'-TACAGGTGACGCCAGGACCTT-5'

This region is transcribed into RNA in a left-to-right direction, and because RNA grows by the addition of successive nucleotides to the 3' end (Figure 1.13), it is the bottom strand that is transcribed. The nucleotide sequence of the RNA is that of the top strand of the DNA, except that U replaces T, so the mRNA for amino acids 1 through 7 is

5'-AUGUCCACUGCGGUCCUGGAA-3'

The codons are read from left to right according to the genetic code shown in Table 1.1. Codon AUG codes for Met (methionine), UCC codes for Ser (serine), and so on. Altogether, the amino acid sequence of this region of the polypeptide is

5'-AUG|UCC|ACU|GCG|GUC|CUG|GAA-3'
Met|Ser|Thr|Ala|Val|Leu|Glu

or, in terms of the single-letter abbreviations,

5'-AUG|UCC|ACU|GCG|GUC|CUG|GAA-3'
M | S | T | A | V | L | E

The full decoding operation for this region of the PH gene is shown in **Figure 1.16**. In this figure, the initiation codon AUG is highlighted because some patients with PKU have a mutation in this particular codon. As might be expected from the fact that methionine is the initiation codon for polypeptide synthesis, cells in patients with this particular mutation fail to produce any of the PAH polypeptide. Mutation and its consequences are considered next.

1.6 Mutation

The term **mutation** refers to any heritable change in a gene (or, more generally, in the genetic material) or to the process by which such a change takes place. One type of mutation results in a change in the sequence of bases in DNA. The change may be simple, such as the substitution of one pair of bases in a duplex molecule for a different pair of bases. For example, a C—G pair in a

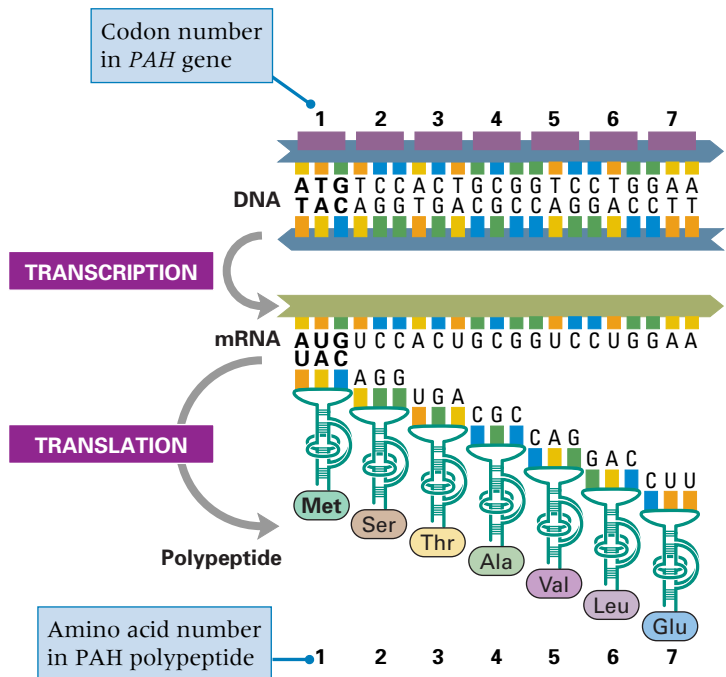


Figure 1.16 The central dogma in action. The DNA that encodes PAH serves as a template for the production of a messenger RNA, and the mRNA serves to specify the sequence of amino acids in the PAH polypeptide chain through interactions with the ribosome and tRNA molecules.

duplex molecule may mutate to T—A, A—T, or G—C. The change in base sequence may also be more complex, such as the deletion or addition of base pairs. These and other types of mutations are considered in Chapter 7. Geneticists also use the term **mutant**, which refers to the result of a mutation. A mutation yields a mutant gene, which in turn produces a mutant mRNA, a mutant protein, and finally a mutant organism that exhibits the effects of the mutation—for example, an inborn error of metabolism.

DNA from patients from all over the world who have phenylketonuria has been studied to determine what types of mutations are responsible for the inborn error. There are a large variety of mutant types. More than 400 different mutations have been described in the gene for PAH. In some cases part of the gene is missing, so the genetic information to make a complete PAH enzyme is absent. In other cases the genetic defect is more subtle, but the result is still either the failure to produce a PAH

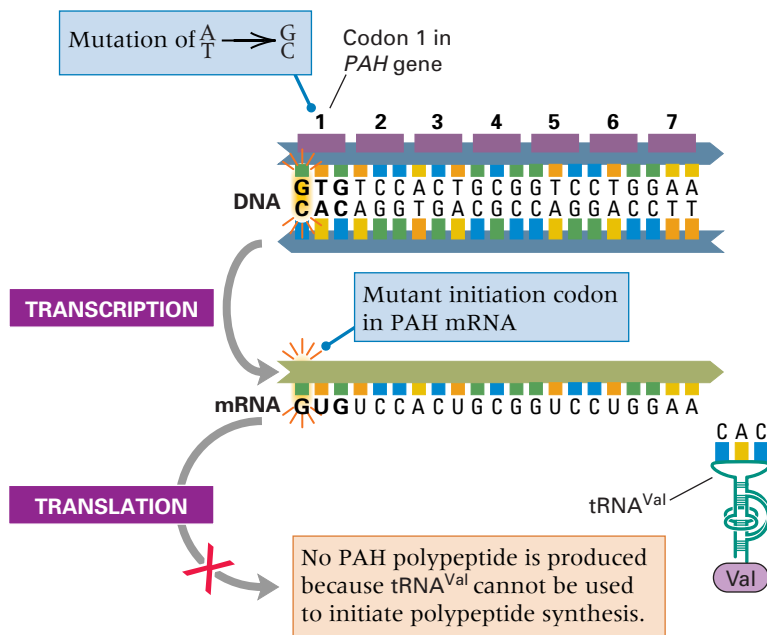


Figure 1.17 The M1V mutant in the *PAH* gene. The methionine codon needed for initiation mutates into a codon for valine. Translation cannot be initiated, and no *PAH* polypeptide is produced.

protein or the production of a *PAH* protein that is inactive. In the mutation shown in **Figure 1.17**, substitution of a G–C base pair for the normal A–T base pair at the very first position in the coding sequence

changes the normal codon AUG (Met) used for the initiation of translation into the codon GUG, which normally specifies valine (Val) and cannot be used as a “start” codon. The result is that translation of the *PAH* mRNA cannot occur, and so no *PAH* polypeptide is made. This mutant is designated M1V because the codon for M (methionine) at amino acid position 1 in the *PAH* polypeptide has been changed to a codon for V (valine). Although the M1V mutant is quite rare worldwide, it is common in some localities, such as Québec Province in Canada.

One *PAH* mutant that is quite common is designated R408W, which means that codon 408 in the *PAH* polypeptide chain has been changed from one coding for arginine (R) to one coding for tryptophan (W). This mutation is one of the four most common among European Caucasians with PKU. The molecular basis of the mutant is shown in **Figure 1.18**. In this case, the first base pair in codon 408 is changed from a C–G base pair into a T–A base pair. The result is that the *PAH* mRNA has a mutant codon at position 408; specifically, it has UGG instead of CGG. Translation does occur in this mutant because everything else about the mRNA is normal, but the result is that the mutant *PAH* carries a tryptophan (Trp) instead of an arginine (Arg) at posi-



The women in the wedding photograph are sisters. Both are homozygous for the same mutant *PAH* gene. The bride is the younger of the two. She was diagnosed just three days after birth and put on the PKU diet soon after. Her older sister, the maid of honor, was diagnosed too late to begin the diet and is mentally retarded. The two-year old pictured in the photo at the right is the daughter of the married couple. They planned the pregnancy: dietary control was strict from conception to delivery to avoid the hazards of excess phenylalanine harming the fetus. Their daughter has passed all developmental milestones with distinction. [Courtesy of Charles R. Scriver.]

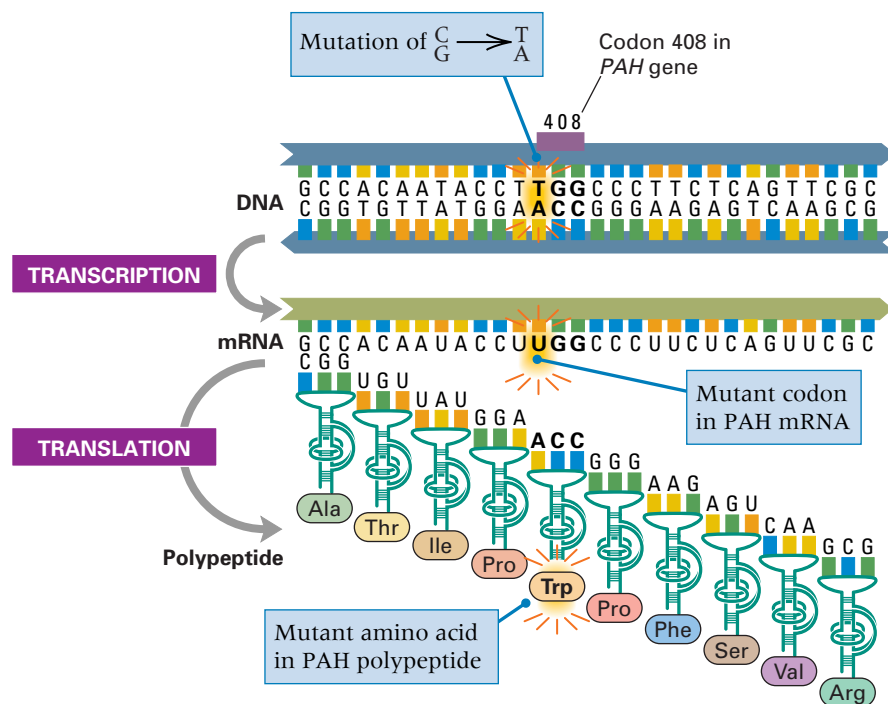


Figure 1.18 The R408W mutant in the *PAH* gene. Codon 408 for arginine (R) is mutated into a codon for tryptophan (W). The result is that position 408 in the mutant PAH polypeptide is occupied by tryptophan rather than by arginine. The mutant protein has no PAH enzyme activity.

tion 408 in the polypeptide chain. The consequence of the seemingly minor change of one amino acid is very drastic. Although the R408W polypeptide is complete, the enzyme has less than 3 percent of the activity of the normal enzyme.

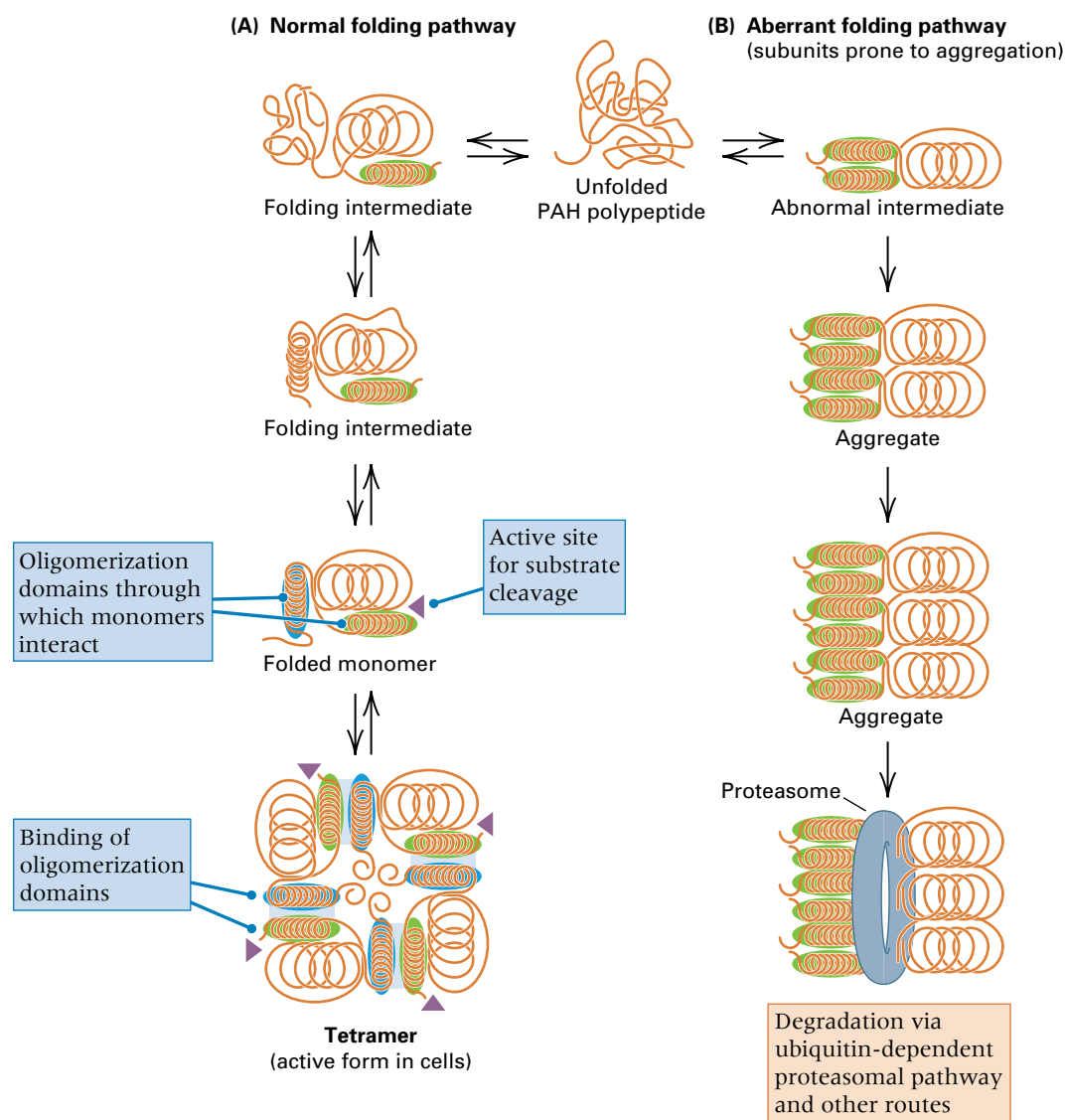
Protein Folding and Stability

More than 400 different mutations in the *PAH* gene have been identified in patients with PKU throughout the world. Many of the mutations affect the level of expression of the gene or processing of the RNA transcript, and some mutations are deletions in which part of the gene is missing. But more than 240 of the mutations are simple amino acid replacements resulting from single nucleotide substitutions in the DNA. Surprisingly, only a minority of amino acid replacements result in a normal amount of PAH protein with a reduced enzyme activity. As a result of most mutations the amount of PAH is reduced, sometimes drastically, and in some other mutations the enzyme activity of the PAH protein that remains is virtually normal. Yet in all these cases the level of expression of the gene, and the amount of mRNA, are within the normal range.

The reason why so many amino acid replacements reduce the amount of protein is that they cause problems in protein fold-

ing, or in the coming together of the protein subunits, or in the stability of the folded protein. **Protein folding** is the complex process by which polypeptide chains attain a stable three-dimensional structure through short-range chemical interactions between nearby amino acids and long-range interactions between amino acids in different parts of the molecule. Folding normally occurs as the polypeptide is being synthesized on the ribosome, and the process is facilitated by a class of proteins known as **chaperones**. During the folding process, the polypeptide chain twists and bends until it achieves a minimum energy state that maximizes the stability of the resulting structure, which is referred to as the **native conformation**. For example, one aspect of protein folding is that hydrophobic amino acids, which have low affinity for water molecules, tend to move toward each other and to form a relatively hydrophobic center, or core, in the native conformation. For a polypeptide of realistic length, there are so many short-range and long-range interactions, and so many possible folded conformations, that even the fastest computers cannot calculate and compare all their energy levels. Computer simulation of protein folding has yielded some insights, but the reliable prediction of protein folding is still a major challenge.

Figure 1.19 Some amino acid replacements perturb the ability of a protein to fold properly. (A) Normal folding in phenylalanine hydroxylase forms the active tetramer. (B) Abnormal folding of a mutant polypeptide chain results in the formation of polypeptide aggregates, which are progressively cleaved into the constituent amino acids through a ubiquitin-dependent proteosomal degradation pathway.



Hypothetical pathways of protein folding in the case of PAH are illustrated in **Figure 1.19**. The normal pathway is shown in part A, including some of the (typically short-lived) intermediates in the folding process. The native conformation of a single PAH polypeptide constitutes the PAH **monomer**. Like many other polypeptides, the PAH monomer contains short regions, called **oligomerization domains** (*oligo* means “a small number”), through which PAH polypeptides undergo stable binding to one another. In the case of PAH, the active form of the PAH enzyme is a **tetramer**, consisting of four identical polypeptide chains held together by interactions between the tetramerization domains. Note that the folding and tetramerization proc-

esses are reversible, so any amino acid replacement that decreases the stability of the tetramer or any of the intermediates will cause more of the PAH polypeptides to fold according to pathway B.

Pathway B in the figure is a misfolding pathway in which the folded monomers are prone to undergo *irreversible* aggregation with each other. These aggregates are targeted for enzymatic breakdown into their constituent amino acids, first by becoming covalently bound with a 76-amino acid polypeptide called *ubiquitin*, which is attached through the activity of several proteins, including a *ubiquitin-conjugating enzyme*. The tagged protein is then degraded by the *proteasome*, which is a large multi-protein complex containing proteins with



Au: Please confirm that caption is OK as set.

Three-dimensional structure of two of the four subunits in the active tetramer of phenylalanine hydroxylase. The chains of amino acids are represented as a sequence of curls, loops, and flat arrows, which represent different types of local structure described in Chapter 11. The oligomerization (in this case, tetramerization) domain is shown in green, and the catalytic domain containing the active site is shown in gold. Compare this with the simplified diagram in Figure 1.19. [Courtesy of R. C. Stevens, T. Flatmark, and H. Erlandsen. From H. Erlandsen, F. Fusetti, A. Martinez, E. Hough, T. Flatmark, and R. C. Stevens. 1997. *Nature Structural Biology* 4: 995.]

ubiquitin-binding, protease, and other activities. The ubiquitin/proteasome pathway is required for the degradation of specific proteins during the cell cycle and development, and it is also used to degrade certain proteins that are intrinsically unstable or that become unfolded in response to stress. In the case of PAH, many amino acid replacements decrease the stability of the protein to such an extent that the protein is routed through the misfolding pathway in Figure 1.19B and targeted for degradation. A destabilizing amino acid replacement can be located at virtually any position along the polypeptide chain.

1.7 Genes and Environment

Inborn errors of metabolism illustrate the general principle that genes code for proteins and that mutant genes code for mutant proteins. In cases such as PKU, mutant proteins cause such a drastic change in metabolism that a severe genetic defect results. But biology is not necessarily destiny. Organisms are also affected by the environment. PKU serves as an example of this principle, because patients who adhere to a diet restricted in the amount of phenylalanine develop mental capacities within the normal range. What is true in this ex-

ample is true in general. Most traits are determined by the interaction of genes and environment.

It is also true that most traits are affected by multiple genes. No one knows how many genes are involved in the development and maturation of the brain and nervous system, but the number must be in the thousands. This number is in addition to the genes that are required in all cells to carry out metabolism and other basic life functions. It is easy to lose sight of the multiplicity of genes when considering extreme examples, such as PKU, in which a single mutation can have such a drastic effect on mental development. The situation is the same as that with any complex machine. An airplane can function if thousands of parts are working together in harmony, but only one defective part, if that part affects a vital system, can bring it down. Likewise, the development and functioning of every trait require a large number of genes working in harmony, but in some cases a single mutant gene can have catastrophic consequences.

In other words, the relationship between a gene and a trait is not necessarily a simple one. The biochemistry of organisms is a complex branching network in which different enzymes may share substrates, yield the same products, or be responsive to the same regulatory elements. The result is

that most visible traits of organisms are the net result of many genes acting together and in combination with environmental factors. PKU affords examples of each of three principles governing these interactions:

1. One gene can affect more than one trait.

Children with extreme forms of PKU often have blond hair and reduced body pigment. This is because the absence of PAH is a metabolic block that prevents conversion of phenylalanine into tyrosine, which is the precursor of the pigment melanin. The relationship between severe mental retardation and decreased pigmentation in PKU makes sense only if one knows the metabolic connections among phenylalanine, tyrosine, and melanin. If these connections were not known, the traits would seem completely unrelated. PKU is not unusual in this regard. Many mutant genes affect multiple traits through their secondary or indirect effects. The various, sometimes seemingly unrelated, effects of a mutant gene are called **pleiotropic effects**, and the phenomenon itself is known as **pleiotropy**.

Figure 1.20 shows a cat with white fur and blue eyes, a pattern of pigmentation that is often (about 40 percent of the time) associated with deafness. Hence deafness can be regarded as a pleiotropic effect of white coat and blue eye color. The developmental basis of this pleiotropy is unknown.



Figure 1.20 Among cats with white fur and blue eyes, about 40 percent are born deaf. We do not know why there is defective hearing nor why it is so often associated with coat and eye color. This form of deafness can be regarded as a pleiotropic effect of white fur and blue eyes.

2. Any trait can be affected by more than one gene.

We discussed this principle earlier in connection with the large number of genes that are required for the normal development and functioning of the brain and nervous system. Among these are genes that affect the function of the blood–brain barrier, which consists of specialized glial cells wrapped around tight capillary walls in the brain, forming an impediment to the passage of most water-soluble molecules from the blood to the brain. The blood–brain barrier therefore affects the extent to which excess free phenylalanine in the blood can enter the brain itself. Because the effectiveness of the blood–brain barrier differs among individuals, PKU patients with very similar levels of blood phenylalanine can have dramatically different levels of cognitive development. This also explains in part why adherence to a controlled-phenylalanine diet is critically important in children but less so in adults; the blood–brain barrier is less well developed in children and is therefore less effective in blocking the excess phenylalanine.

Multiple genes affect even simpler metabolic traits. Phenylalanine breakdown and excretion serve as a convenient example. The metabolic pathway is illustrated in Figure 1.10. Four enzymes in the pathway are indicated, but even more enzymes are involved at the stage labeled “further breakdown.” Because differences in the activity of any of these enzymes can affect the rate at which phenylalanine can be broken down and excreted, all of the enzymes in the pathway are important in determining the amount of excess phenylalanine in the blood of patients with PKU.

3. Most traits are affected by environmental factors as well as by genes.

Here we come back to the low-phenylalanine diet. Children with PKU are not doomed to severe mental deficiency. Their capabilities can be brought into the normal range by dietary treatment. PKU serves as an example of what motivates geneticists to try to discover the molecular basis of inherited disease. The hope is that knowing the metabolic basis of the disease will eventually make it possible to develop methods for clinical intervention through diet, medication, or other treatments that will ameliorate the severity of the disease.

1.8 Evolution: From Genes to Genomes, from Proteins to Proteomes

The pathway for the breakdown and excretion of phenylalanine is by no means unique to human beings. One of the remarkable generalizations to have emerged from molecular genetics is that organisms that are very distinct—for example, plants and animals—share many features in their genetics and biochemistry. These similarities indicate a fundamental “unity of life”:

All creatures on Earth share many features of the genetic apparatus, including genetic information encoded in the sequence of bases in DNA, transcription into RNA, and translation into protein on ribosomes with the use of transfer RNAs. All creatures also share certain characteristics in their biochemistry, including many enzymes and other proteins that are similar in amino acid sequence, three-dimensional structure, and function.

The totality of DNA in a single cell is called the **genome** of the organism. In sexual organisms, the genome is usually regarded as the DNA present in a reproductive cell. The human genome, which is contained in the chromosomes of a sperm or egg, includes approximately 3 billion nucleotide pairs of DNA. The complete set of proteins encoded in the genome is known as the **proteome**. The study of genomes constitutes *genomics*; the study of proteomes constitutes *proteomics*. The fundamental unity of life can be seen in the similarity of proteins in the proteomes of diverse types of organisms. For example, the proteome of the fruit fly *Drosophila melanogaster* includes 13,601 proteins; these can be grouped into 8065 different *families* of proteins that are similar in amino acid sequence. For comparison, the proteome of the nematode worm *Caenorhabditis elegans* contains 18,424 proteins that can be grouped into 9453 families. These two proteomes share about 5000 proteins that are sufficiently similar to be regarded as having a common function. Among the protein families that are common to flies and worms, about 3000 are also found in the proteome of the yeast *Saccharomyces cerevisiae*. Based on these data from these three completely sequenced complex genomes, it seems likely that all organisms whose cells have a nucleus and

chromosomes will turn out to share several thousand protein families. Furthermore, about a thousand of these protein families are shared with organisms as distantly related as bacteria.

The Molecular Unity of Life

Why do organisms share a common set of similar genes and proteins? Because all creatures share a common origin. The process of **evolution** takes place when a population of organisms descended from a common ancestor gradually changes in genetic composition through time. From an evolutionary perspective, the unity of fundamental molecular processes is derived by inheritance from a distant common ancestor in which the molecular mechanisms were already in place.

Not only the unity of life but also many other features of living organisms become comprehensible from an evolutionary perspective. For example, the interposition of an RNA intermediate in the basic flow of

09131_01_1677P

Photocaptionphotoca
ptionphotocaption-
photocaptionphoto-
captionphotocaptionp
hotocaptionphotocap-
tionphotocaptionpho-
tocaptionphotocaption
photocaptionphoto-
captionphotocaption-
photocaptionphotocap
tionphotocaptionpho-
tocaptionphotocap-
tionphotocaption

09131_01_1678P

genetic information from DNA to RNA to protein makes sense if the earliest forms of life used RNA for both genetic information and enzyme catalysis. The importance of the evolutionary perspective in understanding aspects of biology that seem pointless or needlessly complex is summed up in a famous aphorism of the evolutionary biologist Theodosius Dobzhansky: “Nothing in biology makes sense except in the light of evolution.”

One indication of the common ancestry among Earth’s creatures is illustrated in **Figure 1.21**. The tree of relationships was inferred from similarities in nucleotide sequence in an RNA molecule found in the small subunit of the ribosome. Three major kingdoms of organisms are distinguished:

1. **The kingdom Bacteria** This group includes most bacteria and cyanobacteria (formerly called blue-green algae). Cells of these organisms lack a membrane-bounded nucleus and mitochondria, are surrounded by a cell wall, and divide by binary fission.
2. **The kingdom Archaea** This group was initially discovered among microorganisms that produce methane gas or that live in extreme environments, such as hot springs or high salt concentrations. They are widely distributed in more normal environments as well. Like those of bacteria, the cells of archaeans lack internal membranes. DNA sequence analysis indicates that the machinery for DNA replication and transcription in archaeans resembles that of eukaryans, whereas their metabolism strongly

resembles that of bacteria. About half of the genes found in the kingdom Archaea are unique to this group.

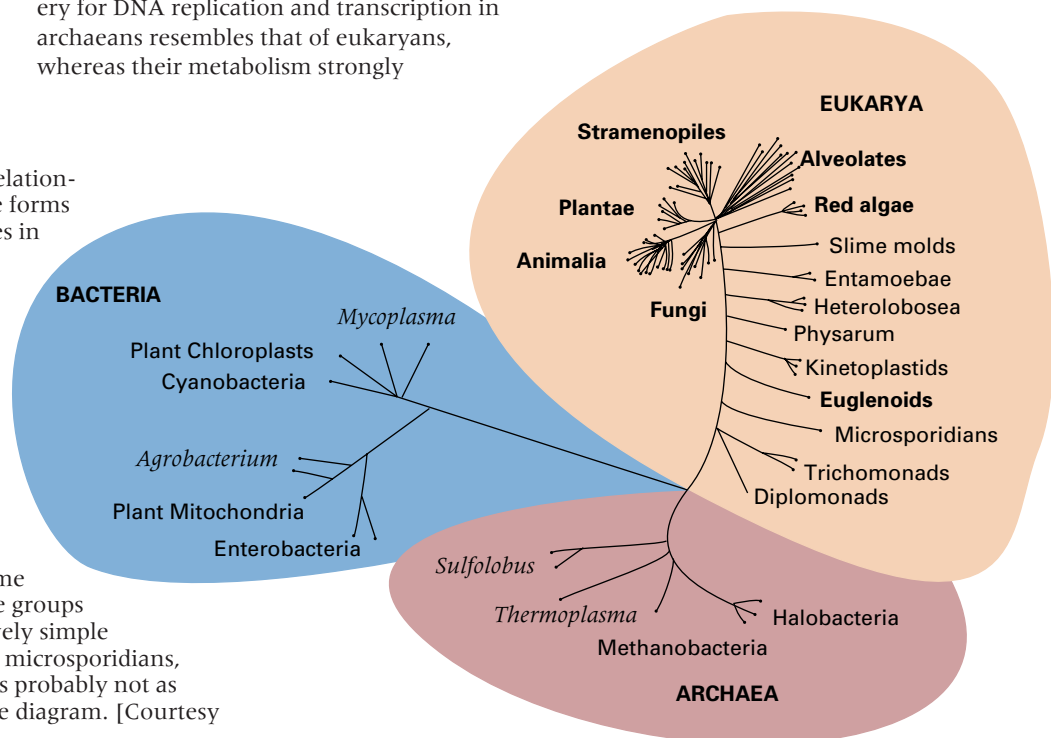
3. **The kingdom Eukarya** This group includes all organisms whose cells contain an elaborate network of internal membranes, a membrane-bounded nucleus, and mitochondria. Their DNA is organized into true chromosomes, and cell division takes place by means of mitosis (discussed in Chapter 4). The eukaryotes include plants and animals as well as fungi and many single-celled organisms, such as amoebae and ciliated protozoa.

The members of the kingdoms Bacteria and Archaea are often grouped together into a larger assemblage called **prokaryotes**, which literally means “before [the evolution of] the nucleus.” This terminology is convenient for designating prokaryotes as a group in contrast with **eukaryotes**, which literally means “good [well-formed] nucleus.”

Natural Selection and Diversity

Although Figure 1.21 illustrates the unity of life, it also illustrates life’s diversity. Frogs are different from fungi, and beetles are different from bacteria. As a human being, it is sobering to consider that complex, multi-cellular organisms are relatively recent ar-

Figure 1.21 Evolutionary relationships among the major life forms as inferred from similarities in nucleotide sequence in an RNA molecule found in the small subunit of the ribosome. The three major kingdoms—Bacteria, Archaea, and Eukarya—are apparent. Plants, animals, and fungi are more closely related to each other than to members of either of the other kingdoms. Among eukaryotes, the time of branching of the diverse groups of undifferentiated, relatively simple organisms (trichomonads, microsporidians, euglenoids, and so forth) is probably not as ancient as suggested by the diagram. [Courtesy of Mitchell L. Sogin.]



rivals to the evolutionary scene of life on Earth. Animals came later still and primates very late indeed. What about human evolution? In the time scale of Earth history, human evolution is a matter of a few million years—barely a snap of the fingers.

If common ancestry is the source of the unity of life, what is the source of its diversity? Because differences among species are inherited, the original source of the differences must be mutation. However, mutations alone are not sufficient to explain why organisms are adapted to living in their environments—why ocean mammals have special adaptations for swimming and diving, why desert mammals have special adaptations that enable them to survive on minimal amounts of water. Mutations are chance events not directed toward any particular adaptive goal, such as longer fur among mammals living in the Arctic. The

process that accounts for adaptation was described by Charles Darwin in his 1859 book *On the Origin of Species*. Darwin proposed that adaptation is the result of **natural selection**: Individual organisms carrying particular mutations or combinations of mutations that enable them to survive or reproduce more effectively in the prevailing environment leave more offspring than other organisms and so contribute their favorable genes disproportionately to future generations. If this process is repeated throughout the course of many generations, the entire species becomes genetically transformed—that is, it evolves—because a gradually increasing proportion of the population inherits the favorable mutations. Mutation, natural selection, and other features of the field of *population genetics* (also called *evolutionary genetics*) are discussed in Chapter 17.

Chapter Summary

Organisms of the same species have some traits (characteristics) in common but may differ from each other in innumerable other traits. Many of the differences between organisms result from genetic differences, the effects of the environment, or both. Genetics is the study of inherited traits, including those influenced in part by the environment. The elements of heredity consist of genes, which are transmitted from parents to offspring in reproduction. Although the sorting of genes in successive generations was first expressed numerically by Mendel, the chemical basis of genes was discovered by Miescher in the form of a weak acid, deoxyribonucleic acid (DNA). However, experimental proof that DNA is the genetic material did not come until about the middle of the twentieth century.

The first convincing evidence of the role of DNA in heredity came from the experiments of Avery, MacLeod, and McCarty, who showed that genetic characteristics in bacteria could be altered from one type to another by treatment with purified DNA. In studies of *Streptococcus pneumoniae*, they transformed mutant cells unable to cause pneumonia into cells that could do so by treating them with pure DNA from disease-causing forms. A second important line of evidence was the Hershey–Chase experiment. Hershey and Chase showed that the T2 bacterial virus injects primarily DNA into the host bacterium (*Escherichia coli*) and that a much higher proportion of parental DNA, compared with parental protein, is found among the progeny phage.

The three-dimensional structure of DNA, proposed in 1953 by Watson and Crick, gave many clues to the manner in which DNA functions as the genetic material. A molecule of DNA consists of two long chains of nu-

cleotide subunits twisted around each other to form a right-handed helix. Each nucleotide subunit contains any one of four bases: A (adenine), T (thymine), G (guanine), or C (cytosine). The bases are paired in the two strands of a DNA molecule; wherever one strand has an A, the partner strand has a T, and wherever one strand has a G, the partner strand has a C. The base pairing means that the two paired strands in a DNA duplex molecule have complementary base sequences along their lengths. The structure of the DNA molecule suggested that genetic information could be coded in DNA in the sequence of bases. Mutations—changes in the genetic material—result from changes in the sequence of bases, such as the substitution of one nucleotide for another or the insertion or deletion of one or more nucleotides. The structure of DNA also suggested a mode of replication: The two strands of the parental DNA molecule separate, and each individual strand serves as a template for the synthesis of a new, complementary strand.

Most genes code for proteins. More precisely stated, most genes specify the sequence of amino acids in a polypeptide chain. The transfer of genetic information from DNA into protein is a multistep process that includes several types of RNA (ribonucleic acid). Structurally, an RNA strand is similar to a DNA strand except that the “backbone” contains a different sugar (ribose instead of deoxyribose) and RNA contains the base uracil (U) instead of thymine (T). Also, RNA is usually present in cells in the form of single, unpaired strands. The initial step in gene expression is transcription, in which a molecule of RNA is synthesized that is complementary in base sequence to whichever DNA strand is being transcribed.

In polypeptide synthesis, which takes place on a ribosome, the base sequence in the RNA transcript is translated in groups of three adjacent bases (codons). The codons are recognized by different types of transfer RNA (tRNA) through base pairing. Each type of tRNA is attached to a particular amino acid, and when a tRNA base-pairs with the proper codon on the ribosome, the growing end of the polypeptide chain is transferred to the amino acid on the tRNA. The table of all codons and the amino acids they specify is called the genetic code. Special codons specify the “start” (AUG, Met) and “stop” (UAA, UAG, and UGA) of polypeptide synthesis. The reason why various types of RNA are an intimate part of transcription and translation is probably that the earliest forms of life used RNA for both genetic information and enzyme catalysis.

A mutation that alters one or more codons in a gene can change the amino acid sequence of the resulting polypeptide chain synthesized in the cell. Often the altered protein is functionally defective, so an inborn error of metabolism results. One of the first inborn errors of metabolism studied was alkaptonuria; it results from the absence of an enzyme for cleaving homogentisic acid, which accumulates and is excreted in the urine, turning black upon oxidation. Phenylketonuria (PKU) is an inborn error of metabolism that affects the same metabolic pathway. The enzyme defect in PKU results in an inability to convert phenylalanine to tyrosine. Phenylalanine accumulation has catastrophic effects on the development of the brain. Children with the disease have severe mental deficits unless they are treated with a special diet low in phenylalanine.

Most visible traits of organisms result from many genes acting together in combination with environmental factors. The relationship between genes and traits is often complex because (1) every gene potentially affects many traits (pleiotropy), (2) every trait is potentially affected by many genes, and (3) many traits are significantly affected by environmental factors as well as by genes.

All living creatures are united by sharing many features of the genetic apparatus (for example, transcription and translation) and many aspects of metabolism. They share many similar genes in the genome and many similar proteins in the proteome. The unity of life results from all life being of common ancestry and provides evidence for evolution. There is also great diversity among living creatures. The three major kingdoms of organisms are the kingdoms Bacteria (whose members lack a membrane-bounded nucleus), Archaea (whose members share features with both bacteria and eukaryotes but form a distinct group), and Eukarya (which includes all “higher” organisms whose cells have a membrane-bounded nucleus that contains DNA organized into discrete chromosomes). Members of the kingdoms Bacteria and Archaea are often collectively called prokaryotes. The ultimate source of diversity among organisms is mutation, but natural selection is the process by which mutations that enhance survival and reproduction are retained and mutations that are harmful are eliminated. Natural selection, first proposed by Darwin, is therefore the primary mechanism by which organisms become progressively better adapted to their environments.

Key Terms

adenine (A)	evolution	polarity (of DNA or RNA)
alkaptonuria	gene	polypeptide chain
amino acid	genetic code	product molecule
Archaea	genetics	prokaryote
Bacteria	genome	protein folding
bacteriophage	guanine (G)	proteome
base (in DNA or RNA)	homogentisic acid	replication
biochemical pathway	inborn error of metabolism	ribonucleic acid (RNA)
block (in a biochemical pathway)	messenger RNA (mRNA)	ribose
central dogma	metabolic pathway	ribosomal RNA (rRNA)
chaperone	metabolite	ribosome
chromosome	monomer	single-stranded DNA
codon	mutant	substrate molecule
colony	mutation	template
complementary base sequence	native conformation	tetramer
cytosine (C)	natural selection	thymine (T)
deoxyribose	nucleotide	transcript
deoxyribonucleic acid (DNA)	oligomerization domain	transcription
double-stranded DNA	phage	transfer RNA (tRNA)
duplex DNA	phenylalanine hydroxylase (PAH)	transformation
enzyme	phenylketonuria (PKU)	translation
Eukarya	pleiotropic effect	uracil (U)
eukaryote	pleiotropy	Watson–Crick base pairing

Review the Basics

- What were the key experiments showing that DNA is the genetic material?
- How did understanding the molecular structure of DNA give clues to its ability to replicate, to code for proteins, and to undergo mutation?
- Why is pairing of complementary bases a key feature of DNA replication?
- What is the process of transcription and in what ways does it differ from DNA replication?
- What three types of RNA participate in protein synthesis, and what is the role of each type of RNA?
- What is the “genetic code,” and how is it relevant to the translation of a polypeptide chain from a molecule of messenger RNA?
- What is an inborn error of metabolism? How did this concept serve as a bridge between genetics and biochemistry?
- How does the “central dogma” explain Garrod’s discovery that nonfunctional enzymes result from mutant genes?
- Explain why many mutant forms of phenylalanine hydroxylase have a simple amino acid replacement, yet the mutant polypeptide chains are absent or present in very small amounts.?
- What is a pleiotropic effect of a gene mutation? Give an example.
- What are some of the major differences in cellular organization among Bacteria, Archaea, and Eukarya?
- What process was Charles Darwin describing when he wrote the following statement? “As more individuals are produced than can possibly survive, there must in every case be a struggle for existence, either one individual with another of the same species, or with the individuals of distinct species, or with the physical conditions of life.”

Guide to Problem Solving

Problem 1 In the human gene for the β chain of hemoglobin (the oxygen-carrying protein in the red blood cells), the first 30 nucleotides in the amino-acid-coding region have the sequence

3'-TACCACGTGGACTGAGGACTCCTCTTCAGA-5'

What is the sequence of the partner strand?

Answer The base pairing between the strands is A with T and G with C, but it is equally important that the strands in a DNA duplex have opposite polarity. The partner strand is therefore oriented with its 5' end at the left, and the base sequence is

5'-ATGGTGCACCTGACTCCTGAGGAGAAGTCT-3'

Problem 2 If the DNA duplex for the β chain of hemoglobin in Problem 1 were transcribed from left to right, deduce the base sequence of the RNA in this coding region.

Answer To deduce the RNA sequence, we must apply three concepts. First, in the transcription of RNA, the base pairing is such that an A, T, G, or C in the DNA template strand is transcribed as U, A, C, or G, respectively, in the RNA strand. Second, the RNA transcript and the DNA template strand have opposite polarity. Third (and critically for this problem), the RNA transcript is always transcribed in the 5'-to-3' direction, so the 5' end of the RNA is the end synthesized first. This being the case, and considering the opposite polarity, the 3' end of the template strand must be transcribed first. Because we are told that transcription takes place from left to right, we can deduce

that the transcribed strand is that given in Problem 1. The RNA transcript therefore has the base sequence

5'-AUGGUGCACCUGACUCCUGAGGAGAAGUCU-3'

Problem 3 Given the RNA sequence coding for part of human β hemoglobin deduced in Problem 2, what is the amino acid sequence in this part of the β polypeptide chain?

Answer The polypeptide chain is translated in successive groups of three nucleotides (each group constituting a codon), starting at the 5' end of the coding sequence and moving in the 5'-to-3' direction. The amino acid corresponding to each codon can be found in the genetic code table. The first ten amino acids in the polypeptide chain are therefore

5'-AUG|GUG|CAC|CUG|ACU|CCU|GAG|GAG|AAG|UCU-3'
Met|Val|His|Leu|Thr|Pro|Glu|Glu|Lys|Ser

Problem 4 A very important mutation in human hemoglobin occurs in the DNA sequence shown in Problem 1. In this mutation, the T at nucleotide position 20 is replaced with an A, where the numbering is from the left in the strand shown at the top in Problem 1. The mutant hemoglobin is called *sickle-cell hemoglobin*, and it is associated with a severe anemia known as *sickle-cell anemia*. Severe as the genetic disease is, the mutant gene is present at relatively high frequency in some human populations because carriers of the gene, who have only a mild anemia, are more resistant to falciparum malaria than are noncarriers. What is the nucleotide sequence of this region of the

GeNETics on the Web will introduce you to some of the most important sites for finding genetic information on the Internet. To explore these sites, visit the Jones and Bartlett home page at

<http://www.jbpub.com/genetics>

For the book *Genetics: Analysis of Genes and Genomes*, choose the link that says **Enter GeNETics on the Web**. You will be presented with a chapter-by-chapter list of highlighted keywords. Select any highlighted keyword and you will be linked to a Web site containing genetic information related to the keyword.

- James D. Watson once said that he and Francis Crick had no doubt that their proposed DNA structure was essentially correct, because the structure was so beautiful it had to be true! At an internet site accessed by the keyword **DNA**, you can view a large collection of different types of models of DNA structure. Some models highlight the sugar-phosphate backbones, others the A-T and G-C base pairs, still others the helical structure of double-stranded DNA.

- How was **PKU** discovered? The story is told by the Norwegian physician Ivar Fölling: "The stage is set in 1934. A mother with two severely mentally retarded children came to

see my father. . . . She had asked many doctors for help, which none had been able to give. But this woman was unusually persistent and would not accept the situation without explanation. She had also noticed that a peculiar smell always clung to her children. . . . This woman was advised to seek help from my father [who held a professorship of nutrient research at the University Hospital in Norway]. He of course had no real hope of being able to help her. But he did not want to reject her, and he agreed to examine the children. On clinical examination he found nothing of importance, except the [severe mental retardation], which was beyond doubt. [Urine analysis] was part of his thorough routine examination. On adding ferric chloride [to normal urine] the color normally stays brownish, [but in the case of these children] a deep green color developed. He had not seen this reaction before . . . He concluded that two mentally retarded children excreted a substance not found in normal urine. But which substance?" Consult this keyword site for more information on how Asbjörn Fölling finally tracked down the cause of the disease.

- With proper dietary control of blood phenylalanine, patients with PKU can develop normally and lead normal lives.

DNA duplex in sickle-cell hemoglobin (both strands) and that of the messenger RNA, and what is the amino acid replacement that results in sickle-cell hemoglobin?

Answer The mutation is already given as a T → A substitution at position 20. The sequence of the DNA duplex is obtained as in Problem 1, that of the RNA as in Problem 2,

and that of the mutant polypeptide chain as in Problem 3, except that at each step there is one nucleotide (or one amino acid) that differs from the nonmutant. The DNA, RNA, and polypeptide in this region of sickle-cell hemoglobin are as follows, where the differences from the non-mutant gene are in red. The amino acid replacement is glutamic acid → valine.

DNA (transcribed strand)	3'-TAC CAC GTG GAC TGA GGA CAC CTC TTC AGA-5'
DNA (nontranscribed strand)	5'-ATG GTG CAC CTG ACT CCT GTG GAG AAG TCT-3'
	↓
RNA coding region	5'-AUG GUG CAC CUG ACU CCU GUG GAG AAG UCU-3'
Polypeptide chain	Met Val His Leu Thr Pro Val Glu Lys Ser

Problem 4 Answer—Nucleotide sequences

Analysis and Applications

1.1 Classify each of the following statements as true or false.

- Each gene is responsible for only one visible trait.
- Every trait is potentially affected by many genes.
- The sequence of nucleotides in a gene specifies the sequence of amino acids in a protein encoded by the gene.

(d) There is one-to-one correspondence between the set of codons in the genetic code and the set of amino acids encoded.

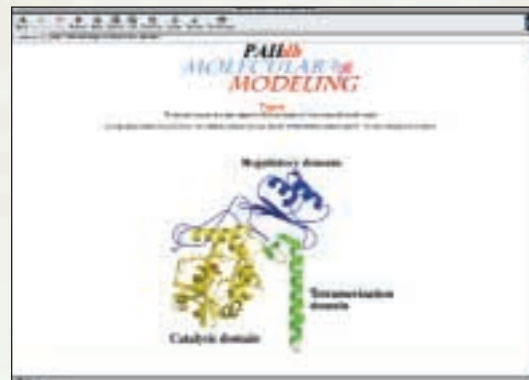
1.2 From their examination of the structure of DNA, what were Watson and Crick able to infer about the probable mechanisms of DNA replication, coding capability, and mutation?

When dietary control is relaxed, however, blood phenylalanine returns to high levels. This situation is extremely dangerous for a developing fetus, resulting in high risk of congenital heart disease, small head size, mental retardation, and slow growth. Affected children are said to have **maternal PKU**. They are affected, not because of their own inability to metabolize phenylalanine, but because of high levels of phenylalanine in their mothers' blood. The risk can be reduced, but not entirely eliminated, if PKU mothers plan their pregnancies and adhere to a strict dietary regimen prior to and during pregnancy. To learn more about this unanticipated consequence of dietary treatment of PKU, log onto the phenylalanine hydroxylase knowledge database at the keyword site.

- Perhaps surprisingly, the history of the bacteriophage T2 that figures so prominently in the experiments of Hershey and Chase is shrouded in mystery. Indeed, the time, place, and source material for the original isolation of phages T2, T4, and T6 (known as the "T-even phages") may never be known with certainty. Use the keyword **T2** to learn what is known about the origin of the bacteriophage and what sleuthing was required to discover what is known.

- The Mutable Site changes frequently. Each new update includes a different site that highlights genetics resources available on the World Wide Web. Select the **Mutable Site** for Chapter 1 and you will be linked automatically.

- The Pic Site showcases some of the most visually appealing genetics sites on the World Wide Web. To visit the genetics Web site pictured below, select the **PIC Site** for Chapter 1.



1.3 What does it mean to say that each strand of a duplex DNA molecule has a polarity? What does it mean to say that the paired strands in a duplex molecule have opposite polarity?

1.4 What is the end result of replication of a duplex DNA molecule?

1.5 What is the role of the messenger RNA in translation? What is the role of the ribosome? What is the role of transfer RNA? Is there more than one type of ribosome? Is there more than one type of transfer RNA?

1.6 What important observation about S and R strains of *Streptococcus pneumoniae* prompted Avery, MacLeod, and McCarty to study this organism?

1.7 In the transformation experiments of Avery, MacLeod, and McCarty, what was the strongest evidence that the substance responsible for the transformation was DNA rather than protein?

1.8 A chemical called phenyl (carbolic acid) destroys proteins but not nucleic acids, and a strong alkali such as sodium hydroxide destroys both proteins and nucleic acids. In the transformation experiments with *Streptococcus pneumoniae*, what result would be expected if the S-strain extract had been treated with phenyl? What would be expected if it had been treated with a strong alkali?

Au: Phenol changed to phenyl two times in prob 1.8. OK?

1.9 What feature of the physical organization of bacteriophage T2 made it suitable for use in the Hershey–Chase experiments?

1.10 Like DNA, molecules of RNA contain large amounts of phosphorus. When Hershey and Chase grew their T2 phage in bacterial cells that had grown in the presence of radioactive phosphorus, the RNA must also have incorporated the labeled phosphorus, and yet the experimental result was not compromised. Why not?

1.11 Although the Hershey–Chase experiments were widely accepted as proof that DNA is the genetic material, the results were not completely conclusive. Why not?

1.12 The DNA extracted from a bacteriophage contains 28 percent A, 28 percent T, 22 percent G, and 22 percent C. What can you conclude about the structure of this DNA molecule?

1.13 The DNA extracted from a bacteriophage consists of 24 percent A, 30 percent T, 20 percent G, and 26 percent C. What is unusual about this DNA? What can you conclude about its structure?

1.14 A double-stranded DNA molecule is separated into its constituent strands, and the strands are separated in an ultracentrifuge. In one of the strands the base composition is 24 percent A, 28 percent T, 22 percent G, and 26

percent C. What is the base composition of the other strand?

1.15 While studying sewage, you discover a new type of bacteriophage that infects *E. coli*. Chemical analysis reveals protein and RNA but no DNA. Is this possible?

1.16 One strand of a DNA duplex has the base sequence 5'-ATCGTATGCACTTTACCCGG-3'. What is the base sequence of the complementary strand?

1.17 A region along one strand of a double-stranded DNA molecule consists of tandem repeats of the trinucleotide 5'-TCG-3', so the sequence in this strand is

5'-TCGTCGTCGTCGTCG ··· -3'

What is the sequence in the other strand?

1.18 A duplex DNA molecule contains a random sequence of the four nucleotides with equal proportions of each. What is the average spacing between consecutive occurrences of the sequence 5'-GGCC-3'? Between consecutive occurrences of the sequence 5'-GAATTC-3'?

1.19 A region along a DNA strand that is transcribed contains no A. What base will be missing in the corresponding region of the RNA?

1.20 The duplex nucleic acid molecule shown here consists of a strand of DNA paired with a complementary strand of RNA. Is the RNA the top or the bottom strand? One of the base pairs is mismatched. Which pair is it?

5'-AUCGGUUAUCCGACUGA-3'
3'-TAGCCAATGTAAGGGTGACT-5'

1.21 The sequence of an RNA transcript that is initially synthesized is 5'-UAGCUAC-3', and successive nucleotides are added to the 3' end. This transcript is produced from a DNA strand with the sequence

3'-AAGTCGCATATCGATGCTAGCGCAACCT-5'

What is the sequence of the RNA transcript when synthesis is complete?

1.22 An RNA molecule folds back upon itself to form a "hairpin" structure held together by a region of base pairing. One segment of the molecule in the paired region has the base sequence 5'-AUACGAUA-3'. What is the base sequence with which this segment is paired?

1.23 A synthetic mRNA molecule consists of the repeating base sequence

5'-UUUUUUUUUUUU ··· -3'

When this molecule is translated *in vitro* using ribosomes, transfer RNAs, and other necessary constituents from *E. coli*, the result is a polypeptide chain consisting of the re-

peating amino acid Phe-Phe-Phe-Phe ···. If you assume that the genetic code is a triplet code, what does this result imply about the codon for phenylalanine (Phe)?

1.24 A synthetic mRNA molecule consisting of the repeating base sequence

5'-UUUUUUUUUUUU ··· -3'

is terminated by the addition, to the right-hand end, of a single nucleotide bearing A. When translated *in vitro*, the resulting polypeptide consists of a repeating sequence of phenylalanines terminated by a single leucine. What does this result imply about the codon for leucine?

1.25 With *in vitro* translation of an RNA into a polypeptide chain, the translation can begin anywhere along the RNA molecule. A synthetic RNA molecule has the sequence

5'-CGCUUACCACAUGUCGCGAACUCG-3'

How many reading frames are possible if this molecule is translated *in vitro*? How many reading frames are possible if this molecule is translated *in vivo*, in which translation starts with the codon AUG?

1.26 You have sequenced both strands of a double-stranded DNA molecule. To inspect the potential amino acid coding content of this molecule, you conceptually transcribe it into RNA and then conceptually translate the RNA into a polypeptide chain. How many reading frames will you have to examine?

1.27 A synthetic mRNA molecule consists of the repeating base sequence 5'-UCUCUCUCUCUCUCUC ··· -3'. When this molecule is translated *in vitro*, the result is a polypeptide chain consisting of the alternating amino acids Ser-Leu-Ser-Leu-Ser-Leu ···. Why do the amino acids alternate? What does this result imply about the codons for serine (Ser) and leucine (Leu)?

1.28 A synthetic mRNA molecule consists of the repeating base sequence 5'-AUCAUCAUCAUCAUC ··· -3'. When this molecule is translated *in vitro*, the result is a mixture of three different polypeptide chains. One consists of repeating isoleucines (Ile-Ile-Ile-Ile ···), another of repeating serines (Ser-Ser-Ser-Ser ···), and the third of repeating histidines (His-His-His-His ···). What does this result imply about the manner in which an mRNA is translated?

1.29 How is it possible for a gene with a mutation in the coding region to encode a polypeptide with the same amino acid sequence as the nonmutant gene?

1.30 A polymer is made that has a random sequence consisting of 75 percent G's and 25 percent U's. Among the amino acids in the polypeptide chains resulting from *in vitro* translation, what is the expected frequency of Trp? of Val? of Phe?

Challenge Problems

1.31 The coding sequence in the messenger RNA for amino acids 1 through 10 of human phenylalanine hydroxylase is

5'-AUGUCCACUGCGGUCCUGGAAAACCCAGGC-3'

- (a) What are the first 10 amino acids?
- (b) What sequence would result from a mutant RNA in which the red A was changed to G?
- (c) What sequence would result from a mutant RNA in which the red C was changed to G?
- (d) What sequence would result from a mutant RNA in which the red U was changed to C?
- (e) What sequence would result from a mutant RNA in which the red G was changed to U?

1.32 A “frameshift” mutation is a mutation in which some number of base pairs, not a multiple of three, is inserted into or deleted from a coding region of DNA. The result is that, at the point of the frameshift mutation, the reading frame of protein translation is shifted with respect to the nonmutant gene. To see the consequence,

consider that the coding sequence in the messenger RNA for the first 10 amino acids in human beta hemoglobin (part of the oxygen carrying protein in the blood) is

5'-AUGGUGCACCUGACUCCUGAGGAGAAGUCU...-3'

- (a) What is the amino acid sequence in this part of the polypeptide chain?
- (b) What would be the consequence of a frameshift mutation resulting in an RNA missing the red U?
- (c) What would be the consequence of a frameshift mutation resulting in an RNA with a G inserted immediately in front of the red U?

1.31 With regard to the wildtype and mutant RNA molecules described in the previous problem, deduce the base sequence in both strands of the corresponding double-stranded DNA for:

- (a) The wildtype sequence
- (b) The single-base deletion
- (c) The single-base insertion

Further Reading

Abedon, S. T. 2000. The murky origin of Snow White and her T-even dwarfs. *Genetics* 155: 481.

Bearn, A. G. 1994. Archibald Edward Garrod, the reluctant geneticist. *Genetics* 137: 1.

Calladine, C. R. 1997. *Understanding DNA: The Molecule and How It Works*. New York: Academic Press.

Ciechanover, A., A. Orian, and A. L. Schwartz. 2000. Ubiquitin-mediated proteolysis: Biological regulation via destruction. *Bioessays* 22: 442.

Doolittle, W. F. 2000. Uprooting the tree of life. *Scientific American*, February.

Gehrig, A., S. R. Schmidt, C. R. Muller, S. Srsen, K. Srsnova, and W. Kress. 1997. Molecular defects in alkaptonuria. *Cytogenetics & Cell Genetics* 76: 14.

Gould, S. J. 1994. The evolution of life on the earth. *Scientific American*, October.

Horowitz, N. H. 1996. The sixtieth anniversary of biochemical genetics. *Genetics* 143: 1.

Judson, H. F. 1996. *The Eighth Day of Creation: The Makers of the Revolution in Biology*. Cold Spring Harbor, NY: Cold Spring Harbor Laboratory Press.

Mirsky, A. 1968. The discovery of DNA. *Scientific American*, June.

Olson, G. J., and C. R. Woese. 1997. Archaeal genomics: An overview. *Cell* 89: 991.

Radman, M., and R. Wagner. 1988. The high fidelity of DNA duplication. *Scientific American*, August.

Rennie, J. 1993. DNA's new twists. *Scientific American*, March.

Scazzocchio, C. 1997. Alkaptonuria: From humans to moulds and back. *Trends in Genetics* 13: 125.

Scriver, C. R., and P. J. Waters. 1999. Monogenic traits are not simple: Lessons from phenylketonuria. *Trends in Genetics* 15: 267.

Stadler, D. 1997. Ultraviolet-induced mutation and the chemical nature of the gene. *Genetics* 145: 863.

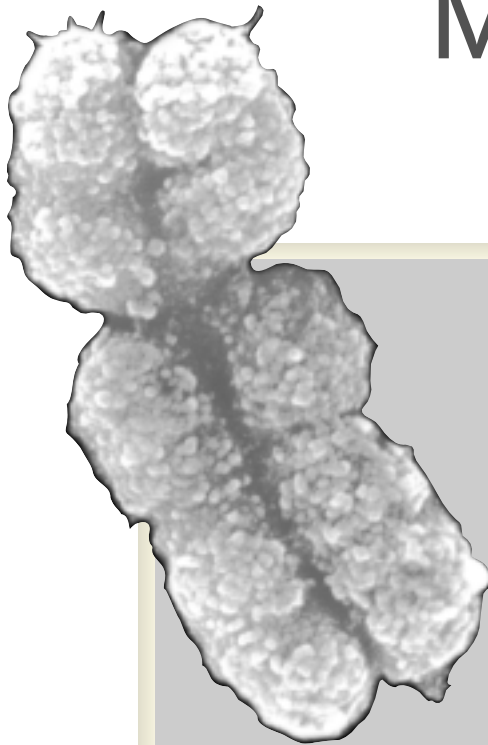
Susman, M. 1995. The Cold Spring Harbor phage course (1945–1970): A 50th anniversary remembrance. *Genetics* 139: 1101.

Watson, J. D. 1968. *The Double Helix*. New York: Atheneum.

Zhang, J. Z. 2000. Protein-length distributions for the three domains of life. *Trends in Genetics* 16: 107.

CHAPTER 2

DNA Structure and DNA Manipulation



CHAPTER OUTLINE

- 2.1 Genomes and Genetic Differences Among Individuals
DNA Markers as Landmarks in Chromosomes
- 2.2 The Molecular Structure of DNA
Polynucleotide Chains
Base Pairing and Base Stacking
Antiparallel Strands
DNA Structure as Related to Function
- 2.3 The Separation and Identification of Genomic DNA Fragments
Restriction Enzymes and Site-Specific DNA Cleavage
Gel Electrophoresis
Nucleic Acid Hybridization
The Southern Blot
- 2.4 Selective Replication of Genomic DNA Fragments
Constraints on DNA Replication:
Primers and 5'-to-3' Strand Elongation
The Polymerase Chain Reaction
- 2.5 The Terminology of Genetic Analysis
- 2.6 Types of DNA Markers Present in Genomic DNA
Single Nucleotide Polymorphisms (SNPs)
Restriction Fragment Length Polymorphisms (RFLPs)
Random Amplified Polymorphic DNA (RAPD)
Amplified Fragment Length Polymorphisms (AFLPs)
Simple Tandem Repeat Polymorphisms (STRPs)
- 2.7 Applications of DNA Markers
Genetic Markers, Genetic Mapping, and "Disease Genes"
Other Uses for DNA Markers

PRINCIPLES

- A DNA strand is a polymer of A, T, G and C deoxyribonucleotides joined 3'-to-5' by phosphodiester bonds.
- The two DNA strands in a duplex are held together by hydrogen bonding between the A–T and G–C base pairs and by base stacking of the paired bases.
- Each type of restriction endonuclease enzyme cleaves double-stranded DNA at a particular sequence of bases usually four or six nucleotides in length.
- The DNA fragments produced by a restriction enzyme can be separated by electrophoresis, isolated, sequenced, and manipulated in other ways.
- Separated strands of DNA or RNA that are complementary in nucleotide sequence can come together (hybridize) spontaneously to form duplexes.
- DNA replication takes place only by elongation of the growing strand in the 5'-to-3' direction through the addition of successive nucleotides to the 3' end.
- In the polymerase chain reaction, short oligonucleotide primers are used in successive cycles of DNA replication to amplify selectively a particular region of a DNA duplex.
- Genetic markers in DNA provide a large number of easily accessed sites in the genome that can be used to identify the chromosomal locations of disease genes, for DNA typing in individual identification, for the genetic improvement of cultivated plants and domesticated animals, and for many other applications.

CONNECTIONS

The Double Helix

James D. Watson and Francis H. C. Crick 1953
A Structure for Deoxyribose Nucleic Acid

Origin of the Human Genetic Linkage Map

David Botstein, Raymond L. White, Mark Skolnick, and Ronald W. Davis 1980
Construction of a Genetic Linkage Map in Man Using Restriction Fragment Length Polymorphisms

In Chapter 1, we reviewed the experimental evidence demonstrating that the genetic material is DNA. We saw how, through the unique structure of the DNA molecule, genetic information can be transcribed and translated into proteins that affect the inherited characteristics of organisms. When a mutant gene encodes a non-functional protein that results in some physical or physiological abnormality—for example, an “inborn error of metabolism”—the expression of that abnormality can be used to trace the transmission of the mutant gene from one generation to the next in a pedigree (family history). As a consequence, until recently, the first step in genetic analysis was the identification of organisms with such abnormal traits, such as peas with wrinkled seeds instead of round seeds and fruit flies with white eyes instead of red. These traits were studied by means of controlled crosses so that the parentage of each individual could be traced. Large-scale genetic studies were typically limited to one of a small number of model organisms especially favorable for isolating and identifying mutant genes, such as the budding yeast *Saccharomyces cerevisiae*, the nematode worm *Caenorhabditis elegans*, or the fruit fly *Drosophila melanogaster*.

Since the mid-1970s, studies in genetics have undergone a revolution based on the use of increasingly sophisticated ways to isolate and identify specific fragments of DNA. The culmination of these techniques was large-scale genomic sequencing—the ability to determine the correct sequence of the base pairs that make up the DNA in an entire genome and to identify the sequences associated with genes. Because many of the model organisms used in genetics have relatively small genomes, these sequences were completed first, in the late 1990s (Figure 2.1). The techniques used to sequence these simpler genomes were then scaled up to sequence the human genome. The initial “rough draft” of the human genome was announced in June 2000; this represents an important milestone in the Human Genome Project, whose goals in-

09131_01_1651A

Figure 2.1 Timeline of large-scale genomic DNA sequencing.

clude determining the sequence, and identifying the function, of all human genes.

2.1 Genomes and Genetic Differences Among Individuals

The numbers associated with the genome of even a simple organism can be intimidating. The sequenced genomes of *D. melanogaster* and *C. elegans*, both approximately 100 million base pairs in length, encode 13,601 proteins and 18,424 proteins, respectively. The human genome is considerably larger. As found in a human reproductive cell, the human genome consists of 3 billion base pairs organized into 23 distinct chromosomes (each chromosome contains a single molecule of duplex DNA). A typical chromosome can contain several hundred to several thousand genes, arranged in linear order along the DNA molecule present in the chromosome. The sequences that make up the protein-coding part of these genes actually account for only about 4 percent of the entire genome. The other 96 percent of the sequences do not code for proteins. Some noncoding sequences are genetic “chaff” that gets separated from the protein-coding “wheat” when genes are transcribed and the RNA transcript is processed into messenger RNA. Other noncoding sequences are relatively short sequences that are found in hundreds or thousands of copies scattered throughout the genome. Still other noncoding sequences are remnants of genes called *pseudogenes*. As might be expected, identifying the protein-coding genes from among the large background of noncoding DNA in the human genome is a challenge in itself.

Geneticists often speak of the nucleotide sequence of “the” human genome because 99.9 percent of the DNA sequences in any two individuals are the same. This is our evolutionary legacy; it contains the genetic information that makes us human beings. In reality, however, there are many different human genomes. Geneticists have great interest in the 0.1 percent of the human DNA sequence—3 million base pairs—that differs from one genome to the next, because these differences include the muta-

tions that are responsible for genetic diseases such as phenylketonuria and other inborn errors of metabolism, as well as the mutations that increase the risk of more complex diseases such as heart disease, breast cancer, and diabetes.

Fortunately, only a small proportion of all differences in DNA sequence are associated with disease. Some of the others are associated with inherited differences in height, weight, hair color, eye color, facial features, and other traits. Most of the genetic differences between people are completely harmless. Many have no detectable effects on appearance or health. Such differences can be studied only through direct examination of the DNA itself. These harmless differences are nevertheless important, because they serve as genetic markers.

DNA Markers as Landmarks in Chromosomes

In genetics, a **genetic marker** is any difference in DNA, no matter how it is detected, whose pattern of transmission from generation to generation can be tracked. Each individual who carries the marker also carries a length of chromosome on either side of it, so it *marks* a particular region of the genome. A mutant gene, or some portion of a mutant gene, can serve as a genetic marker. In the “classical” approach to genetics, it is the outward expression of a gene (or lack of expression) that forms the basis of genetic analysis. For example, a mutation causing wrinkled peas is a genetic marker, which can be identified through its effects on pea shape. In modern genetic analysis, *any* difference in DNA sequence between two individuals can serve as a genetic marker. And although these genetic markers are often harmless in themselves, they allow the positions of disease genes to be located and their DNA isolated, identified, and studied.

Genetic markers that are detected by direct analysis of the DNA are often called **DNA markers**. DNA markers are important in genetics because they serve as landmarks in long DNA molecules, such as those found in chromosomes, which allow genetic differences among individuals to be

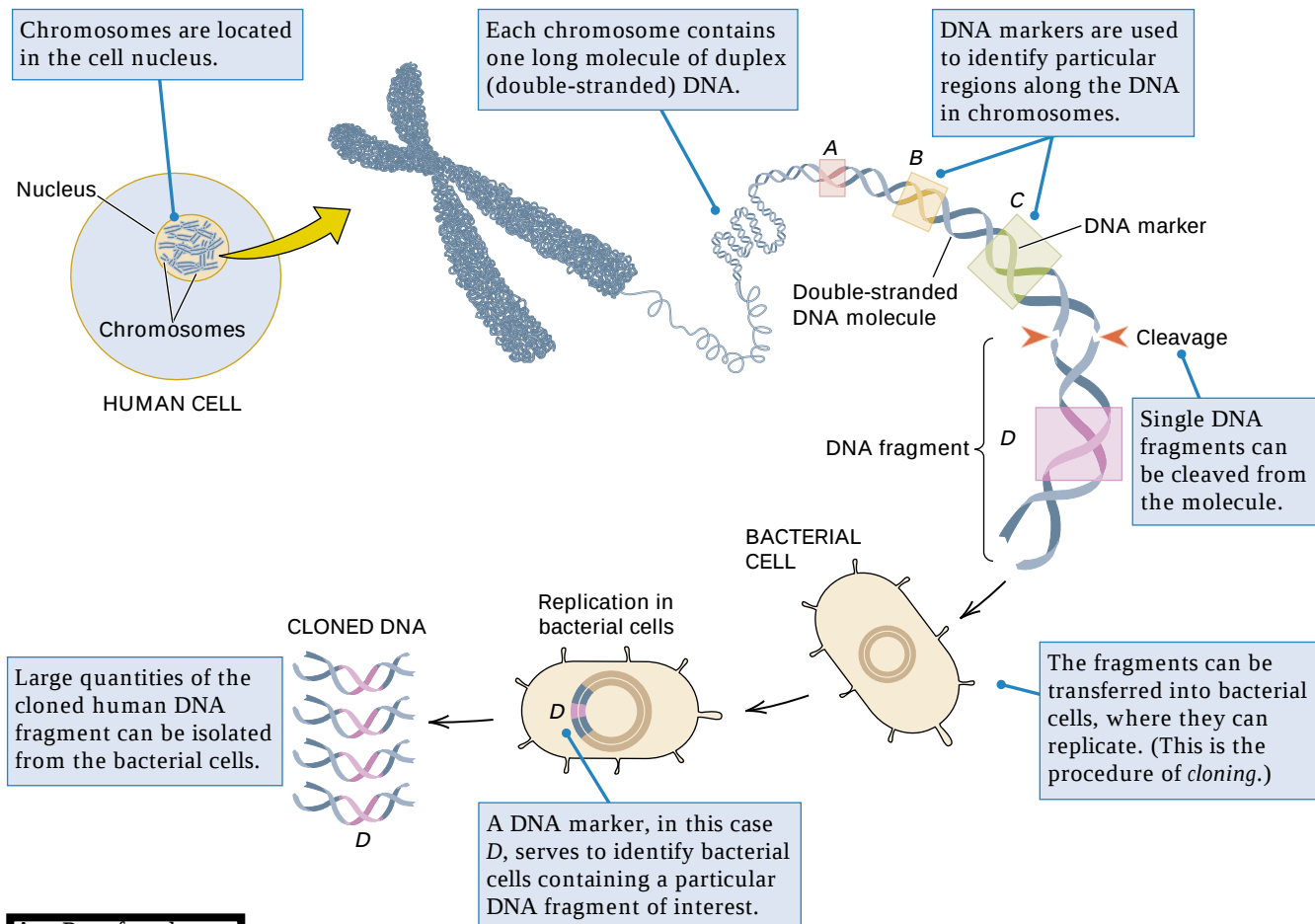
Au: In 1st para in 1st column, proof-reader asks if “hundreds or thousands of copies” is correct? Or should phrase be “hundreds of thousands of copies”?

tracked. They are like signposts along a highway. Using DNA markers as landmarks, the geneticist can identify the positions of normal genes, mutant genes, breaks in chromosomes, and other features important in genetic analysis (Figure 2.2).

The detection of DNA markers usually requires that the **genomic DNA** (the total DNA extracted from cells of an organism) be fragmented into pieces of manageable size (usually a few thousand nucleotide pairs) that can be manipulated in laboratory experiments. In the following sections, we examine some of the principal ways in which DNA is manipulated to reveal genetic differences among individuals, whether or not these differences find out-

ward expression. Use of these methods broadens the scope of genetics, making it possible to carry out genetic analysis in *any* organism. This means that detailed genetic analysis is no longer restricted to human beings, domesticated animals, cultivated plants, and the relatively small number of model organisms favorable for genetic studies. Direct study of DNA eliminates the need for prior identification of genetic differences between individuals; it even eliminates the need for controlled crosses. The methods of molecular analysis discussed in this chapter have transformed genetics:

The manipulation of DNA is the basic experimental operation in modern genetics.



Au: Proofreader asks that you confirm Chapter 13 reference in Fig 2.2 caption is correct.

Figure 2.2 DNA markers serve as landmarks that identify physical positions along a DNA molecule, such as DNA from a chromosome. As shown at the right, a DNA marker can also be used to identify bacterial cells into which a particular fragment of DNA has been introduced. The procedure of DNA cloning is not quite as simple as indicated here; it is discussed further in Chapter 13.

These methods are the principal techniques used in virtually every modern genetics laboratory.

2.2 The Molecular Structure of DNA

Modern experimental methods for the manipulation and analysis of DNA grew out of a detailed understanding of its molecular structure and replication. Therefore, to understand these methods, one needs to know something about the molecular structure of DNA. We saw in Chapter 1 that DNA is a helix of two paired, complementary strands, each composed of an ordered string of *nucleotides*, each bearing one of the bases A (adenine), T (thymine), G (guanine), or cytosine (C). Watson–Crick base pairing between A and T and between G and C in the complementary strands holds the strands together. The complementary strands also hold the key to replication, because each strand can serve as a template for the synthesis of a new complementary strand. We will now take a closer look at DNA structure and at the key features of its replication.

Polynucleotide Chains

In terms of biochemistry, a DNA strand is a polymer—a large molecule built from repeating units. The units in DNA are composed of 2'-deoxyribose (a five-carbon su-



Photocaptionphotocaptionphotocaptionphotocaptionphotocaptionphotocaptionphotocaptionphotocaptionphotocaptionphotocaption

gar), phosphoric acid, and the four nitrogen-containing bases denoted A, T, G, and C. The chemical structures of the bases are shown in Figure 2.3. Note that two of the bases have a double-ring structure; these are called **purines**. The other two bases have a single-ring structure; these are called **pyrimidines**.

- The purine bases are adenine (A) and guanine (G).
- The pyrimidine bases are thymine (T) and cytosine (C).

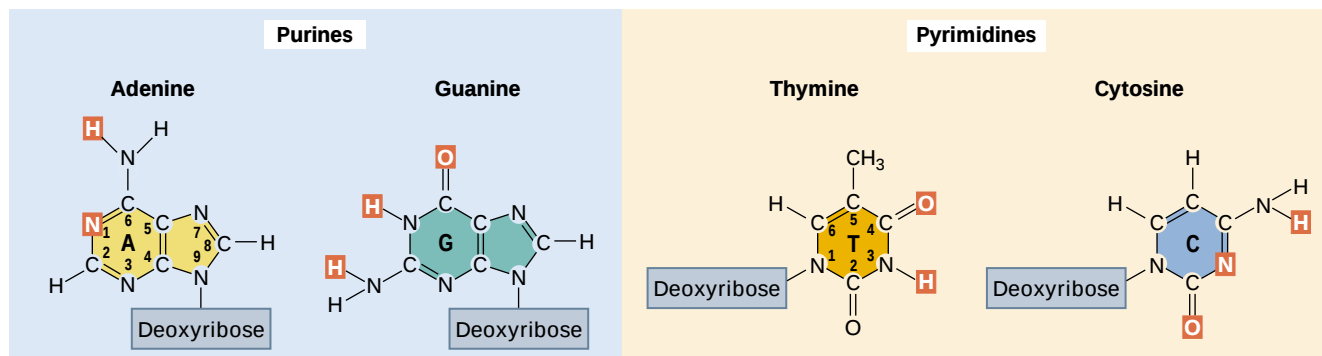


Figure 2.3 Chemical structures of the four nitrogen-containing bases in DNA: adenine, thymine, guanine, and cytosine. The nitrogen atom linked to the deoxyribose sugar is indicated. The atoms shown in red participate in hydrogen bonding between the DNA base pairs.

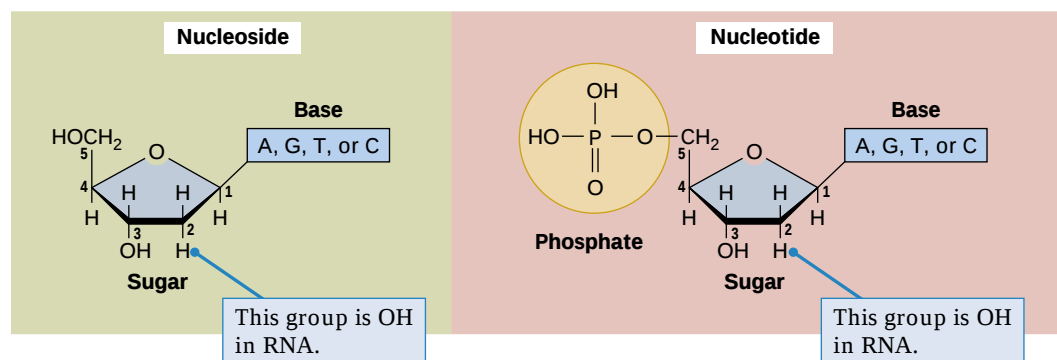


Figure 2.4 A typical nucleotide, showing the three major components (phosphate, sugar, and base), the difference between DNA and RNA, and the distinction between a nucleoside (no phosphate group) and a nucleotide (with phosphate). Nucleotides are monophosphates (with one phosphate group). Nucleoside diphosphates contain two phosphate groups, and nucleoside triphosphates contain three.

In DNA, each base is chemically linked to one molecule of the sugar deoxyribose, forming a compound called a **nucleoside**. When a phosphate group is also attached to the sugar, the nucleoside becomes a **nucleotide** (Figure 2.4). Thus a nucleotide is a nucleoside plus a phosphate. In the conventional numbering of the carbon atoms in the sugar in Figure 2.4, the carbon atom to which the base is attached is the 1' carbon. (The atoms in the sugar are given primed numbers to distinguish them from atoms in the bases.) The nomenclature of the nucleoside and nucleotide derivatives of the DNA bases is summarized

in Table 2.1. Most of these terms are not needed in this book; they are included because they are likely to be encountered in further reading.

In nucleic acids, such as DNA and RNA, the nucleotides are joined to form a **polynucleotide chain**, in which the phosphate attached to the 5' carbon of one sugar is linked to the hydroxyl group attached to the 3' carbon of the next sugar in line (Figure 2.5). The chemical bonds by which the sugar components of adjacent nucleotides are linked through the phosphate groups are called **phosphodiester bonds**. The 5'–3'–5'–3' orientation of these linkages

Au: Query from MH—Why no primes shown here (shown in Fig 2.5)?

Table 2.1 DNA nomenclature		
Base	Nucleoside	Nucleotide
Adenine (A)	Deoxyadenosine	Deoxyadenosine-5' monophosphate (dAMP) diphosphate (dADP) triphosphate (dATP)
Guanine (G)	Deoxyguanosine	Deoxyguanosine-5' monophosphate (dGMP) diphosphate (dGDP) triphosphate (dGTP)
Thymine (T)	Deoxythymidine	Deoxythymidine-5' monophosphate (dTMP) diphosphate (dTDP) triphosphate (dTTP)
Cytosine (C)	Deoxycytidine	Deoxycytidine-5' monophosphate (dCMP) diphosphate (dCDP) triphosphate (dCTP)

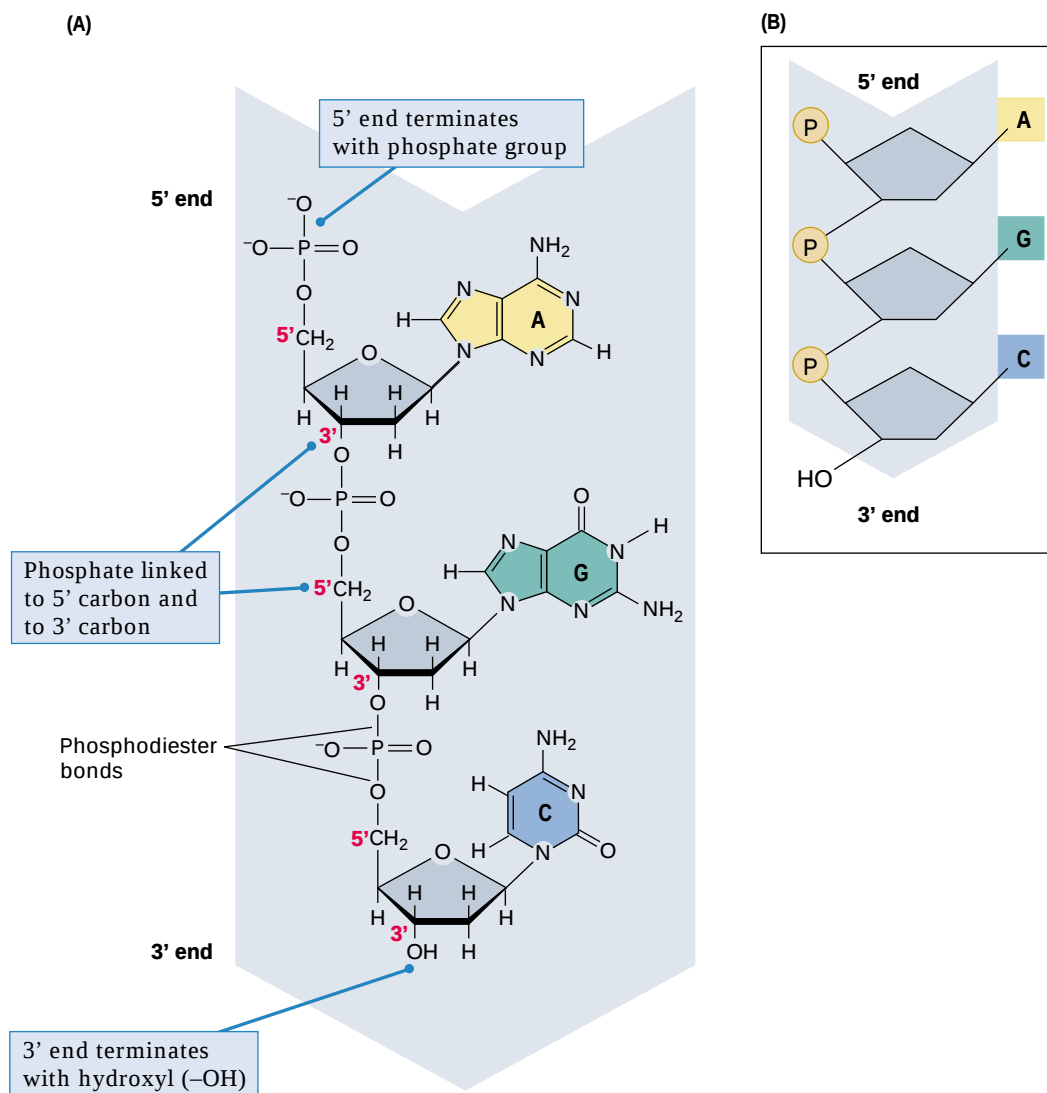


Figure 2.5 Three nucleotides at the 5' end of a single polynucleotide strand. (A) The chemical structure of the sugar-phosphate linkages, showing the 5'-to-3' orientation of the strand (the red numbers are those assigned to the carbon atoms). (B) A common schematic way to depict a polynucleotide strand.

continues throughout the chain, which typically consists of millions of nucleotides. Note that the terminal groups of each polynucleotide chain are a **5'-phosphate (5'-P) group** at one end and a **3'-hydroxyl (3'-OH) group** at the other. The asymmetry of the ends of a DNA strand implies that each strand has a **polarity** determined by which end bears the 5' phosphate and which end bears the 3' hydroxyl.

A few years before Watson and Crick proposed their essentially correct three-dimensional structure of DNA as a double helix, Erwin Chargaff developed a chemical

technique to measure the amount of each base present in DNA. As we describe this technique, we will let the molar concentration of any base be represented by the symbol for the base in square brackets; for example, [A] denotes the molar concentration of adenine. Chargaff used his technique to measure the [A], [T], [G], and [C] content of the DNA from a variety of sources. He found that the **base composition** of the DNA, defined as the **percent G + C**, differs among species but is constant in all cells of an organism and within a species. Data on the base composition of

DNA from a variety of organisms are given in Table 2.2.

Chargaff also observed certain regular relationships among the molar concentrations of the different bases. These relationships are now called **Chargaff's rules**:

- The amount of adenine equals that of thymine: $[A] = [T]$.
- The amount of guanine equals that of cytosine: $[G] = [C]$.
- The amount of the purine bases equals that of the pyrimidine bases:

$$[A] + [G] = [T] + [C].$$

Although the chemical basis of these observations was not known at the time, one of the appealing features of the Watson–Crick structure of paired complementary strands was that it explained Chargaff's rules. Because A is always paired with T in double-stranded DNA, it must follow that $[A] = [T]$. Similarly, because G is paired with C, we know that $[G] = [C]$. The third rule follows

by addition of the other two: $[A] + [G] = [T] + [C]$. In the next section, we examine the molecular basis of base pairing in more detail.

Base Pairing and Base Stacking

In the three-dimensional structure of the DNA molecule proposed in 1953 by Watson and Crick, the molecule consists of two polynucleotide chains twisted around one another to form a double-stranded helix in which adenine and thymine, and guanine and cytosine, are paired in opposite strands (Figure 2.6). In the standard structure, which is called the **B form of DNA**, each chain makes one complete turn every 34 Å. The helix is right-handed, which means that as one looks down the barrel, each chain follows a clockwise path as it progresses. The bases are spaced at 3.4 Å, so there are ten bases per helical turn in each strand and ten base pairs per turn of the double helix.

Table 2.2 Base composition of DNA from different organisms					
Organism	Base (and percentage of total bases)				Base composition (percent G + C)
	Adenine	Thymine	Guanine	Cytosine	
Bacteriophage T7	26.0	26.0	24.0	24.0	48.0
Bacteria					
<i>Clostridium perfringens</i>	36.9	36.3	14.0	12.8	26.8
<i>Streptococcus pneumoniae</i>	30.2	29.5	21.6	18.7	40.3
<i>Escherichia coli</i>	24.7	23.6	26.0	25.7	51.7
<i>Sarcina lutea</i>	13.4	12.4	37.1	37.1	74.2
Fungi					
<i>Saccharomyces cerevisiae</i>	31.7	32.6	18.3	17.4	35.7
<i>Neurospora crassa</i>	23.0	22.3	27.1	27.6	54.7
Higher plants					
Wheat	27.3	27.2	22.7	22.8*	45.5
Maize	26.8	27.2	22.8	23.2*	46.0
Animals					
<i>Drosophila melanogaster</i>	30.8	29.4	19.6	20.2	39.8
Pig	29.4	29.6	20.5	20.5	41.0
Salmon	29.7	29.1	20.8	20.4	41.2
Human being	29.8	31.8	20.2	18.2	38.4

*Includes one-fourth 5-methylcytosine, a modified form of cytosine found in most plants more complex than algae and in many animals

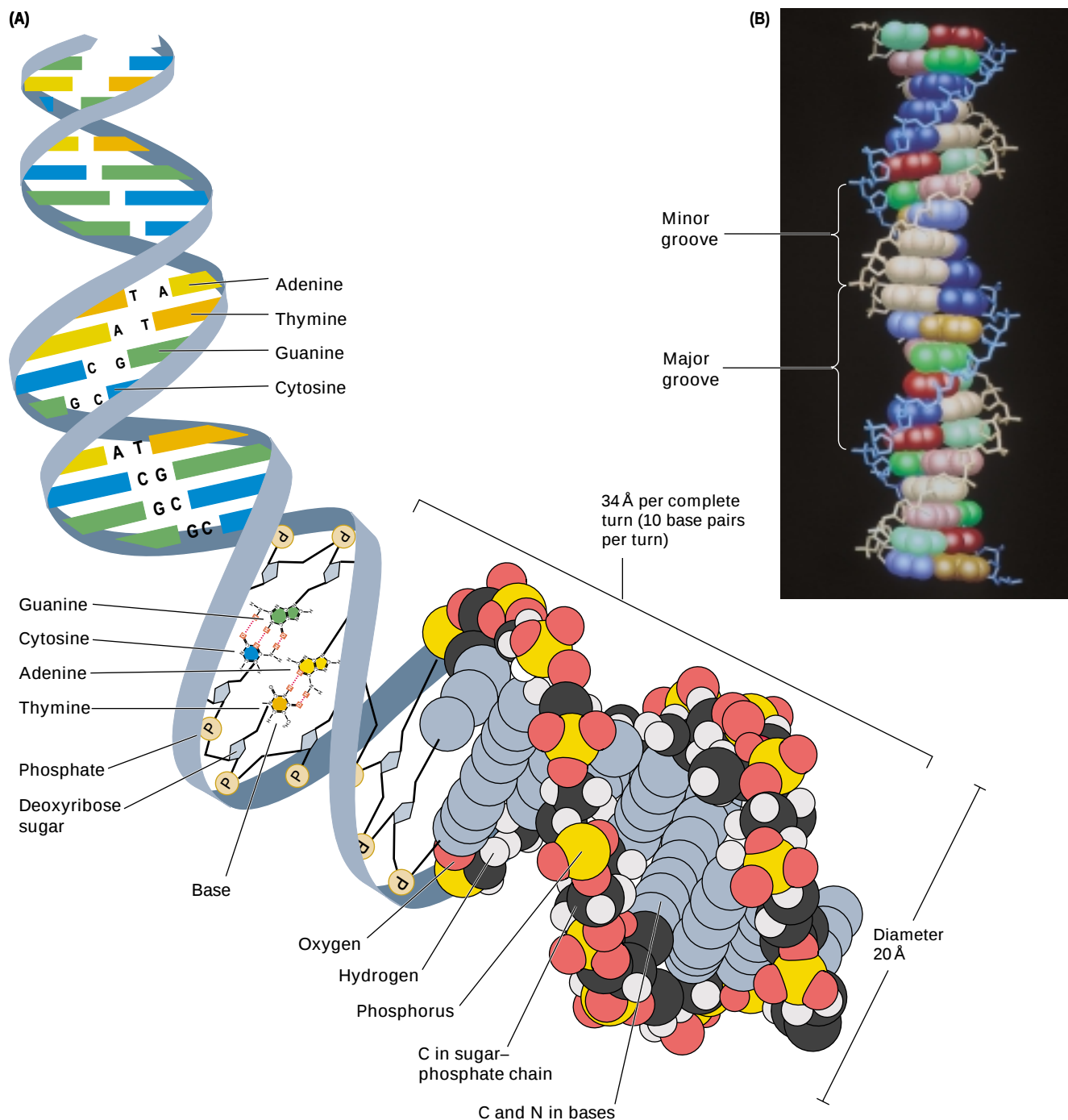


Figure 2.6 Two representations of DNA, illustrating the three-dimensional structure of the double helix. (A) In a ribbon diagram, the sugar-phosphate backbones are depicted as bands, with horizontal lines used to represent the base pairs. (B) A computer model of the B form of a DNA molecule. The stick figures are the sugar-phosphate chains winding around outside the stacked base pairs, forming a major groove and a minor groove. The color coding for the base pairs is as follows: A, red or pink; T, dark green or light green; G, dark brown or beige; C, dark blue or light blue. The bases depicted in dark colors are those attached to the blue sugar-phosphate backbone; the bases depicted in light colors are attached to the beige backbone. [B, courtesy of Antony M. Dean.]

Au: Term in Key Terms list is “hydrophobic interaction” rather than “hydrophic.” Would you like to change term in list? Or change text here?

The strands feature **base pairing**, in which each base is paired to a complementary base in the other strand by hydrogen bonds. (A **hydrogen bond** is a weak bond in which two participating atoms share a hydrogen atom between them.) The hydrogen bonds provide one type of force holding the strands together. In Watson–Crick base pairing, adenine (A) pairs with thymine (T), and guanine (G) pairs with cytosine (C). The hydrogen bonds that form in the adenine–thymine base pair and in the guanine–cytosine pair are illustrated in Figure 2.7. Note that an A–T pair (Figure 2.7A and B) has two hydrogen bonds and that a G–C pair (Figure 2.7C and D) has three hydrogen bonds. This means that the hydrogen bonding between G and C is stronger in the sense that it requires more energy to break; for example, the amount of heat required to separate the paired strands in a DNA duplex increases with the

percent of G + C. Because nothing restricts the sequence of bases in a single strand, any sequence could be present along one strand. This explains Chargaff’s observation that DNA from different organisms may differ in base composition. However, because the strands in duplex DNA are complementary, Chargaff’s rules of $[A] = [T]$ and $[G] = [C]$ are true whatever the base composition.

In the B form of DNA, the paired bases are planar, parallel to one another, and perpendicular to the long axis of the double helix. This feature of double-stranded DNA is known as **base stacking**. The upper and lower faces of each nitrogenous base are relatively flat and nonpolar (uncharged). These surfaces are said to be **hydrophobic** because they bind poorly to water molecules, which are very polar. (The polarity refers to the asymmetrical distribution of charge across the V-shaped water molecule;

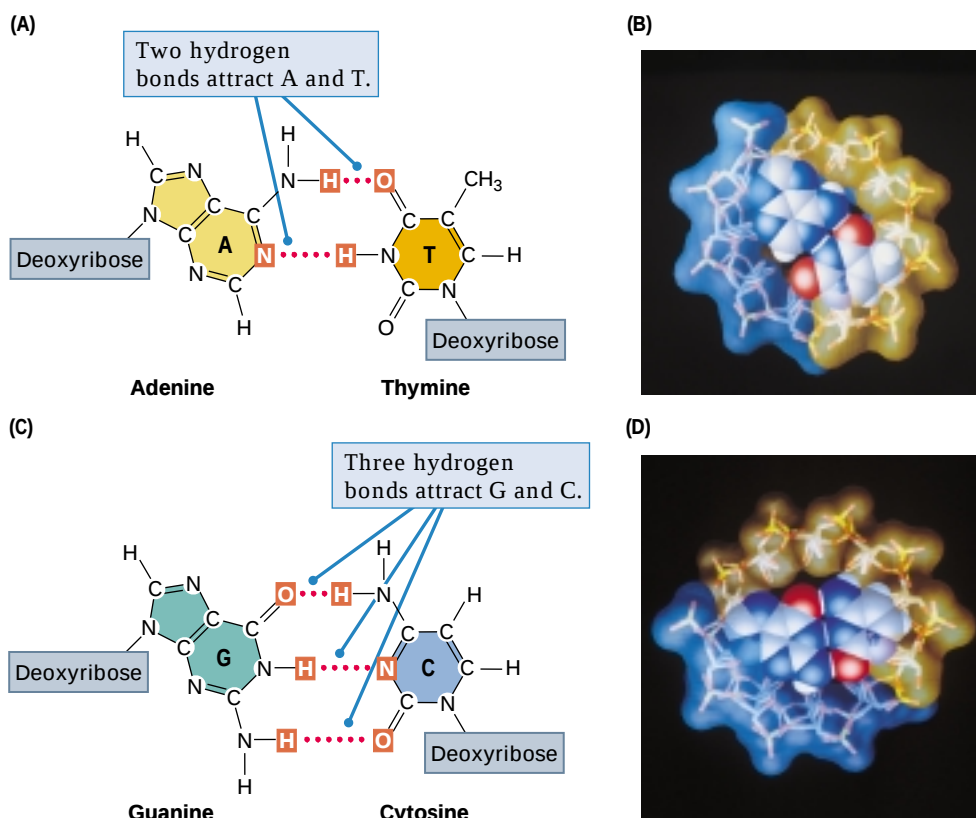


Figure 2.7 Normal base pairs in DNA. On the left, the hydrogen bonds (dotted lines) with the joined atoms are shown in red. (A and B) A–T base pairing. (C and D) G–C base pairing. In the space-filling models (B and D), the colors are as follows: C, gray; N, blue; O, red; and H (shown in the bases only), white. Each hydrogen bond is depicted as a white disk squeezed between the atoms sharing the hydrogen. The stick figures on the outside represent the backbones winding around the stacked base pairs. [Space-filling models courtesy of Antony M. Dean.]

the oxygen at the base of the V tends to be quite negative, whereas the hydrogens at the tips are quite positive). Owing to their repulsion of water molecules, the paired nitrogenous bases tend to stack on top of one another in such a way as to exclude the maximum amount of water from the interior of the double helix. Hence a double-stranded DNA molecule has a hydrophobic core composed of stacked bases, and it is the energy of base stacking that provides double-stranded DNA with much of its chemical stability.

When discussing a DNA molecule, molecular biologists frequently refer to the individual strands as single strands or as single-stranded DNA; they refer to the double helix as double-stranded DNA or as a *duplex* molecule. The two grooves spiraling along outside of the double helix are not symmetrical; one groove, called the **major groove**, is larger than the other, which is called the **minor groove**. Proteins that interact with double-stranded DNA often have regions that make contact with the base pairs by fitting into the major groove, into the minor groove, or into both grooves (Figure 2.6B).

Antiparallel Strands

Each backbone in a double helix consists of deoxyribose sugars alternating with phosphate groups that link the 3' carbon atom of one sugar to the 5' carbon atom of the next in line (Figure 2.5). The two polynucleotide strands of the double helix have opposite polarity in the sense that the 5' end of one strand is paired with the 3' end of the other strand. Strands with such an arrangement are said to be **antiparallel**. One implication of antiparallel strands in duplex DNA is that in each pair of bases, one base is attached to a sugar that lies above the plane of pairing, and the other base is attached to a sugar that lies below the plane of pairing. Another implication is that each terminus of the double helix possesses one 5'-P group (on one strand) and one 3'-OH group (on the other strand), as shown in Figure 2.8.

The diagram of the DNA duplex in Figure 2.6 is static and so somewhat misleading. DNA is a dynamic molecule, constantly in motion. In some regions, the strands can separate briefly and then come together again in the same conformation or in a dif-

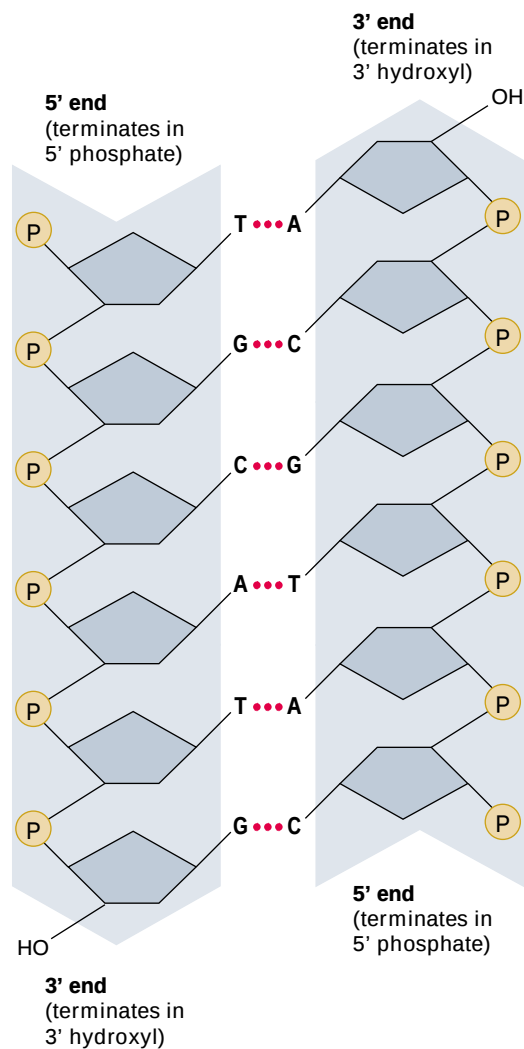


Figure 2.8 A segment of a DNA molecule, showing the antiparallel orientation of the complementary strands. The overlying blue arrows indicate the 5'-to-3' direction of each strand. The phosphates (P) join the 3' carbon atom of one deoxyribose (horizontal line) to the 5' carbon atom of the adjacent deoxyribose.

ferent one. The right-handed double helix in Figure 2.6 is the standard B form, but depending on conditions, DNA can actually form more than 20 slightly different variants of a right-handed helix, and some regions can even form helices in which the strands twist to the left (called the **Z form of DNA**). If there are complementary stretches of nucleotides in the same strand, then a single strand, separated from its partner, can fold back upon itself like a hairpin. Even triple helices consisting of three strands can form in regions of DNA that contain suitable base sequences.

James D. Watson and
Francis H. C. Crick 1953

Cavendish Laboratory,
Cambridge, England

A Structure for Deoxyribose Nucleic Acid

This is one of the watershed papers of twentieth-century biology. After its publication, nothing in genetics was the same. Everything that was known, and everything still to be discovered, would now need to be interpreted in terms of the structure and function of DNA. The importance of the paper was recognized immediately, in no small part because of its lucid and concise description of the structure. Watson and Crick benefited tremendously in knowing that their structure was consistent with the unpublished structural studies of Maurice Wilkins and Rosalind Franklin. The same issue of Nature that included the Watson and Crick paper also included, back to back, a paper from the Wilkins group and one from the Franklin group detailing their data and the consistency of their data with the proposed structure. It has been said that Franklin was poised a mere two half-steps from making the discovery herself, alone. In any event, Watson and Crick and Wilkins were awarded the 1962 Nobel Prize for their discovery of DNA structure. Rosalind Franklin, tragically, died of cancer in 1958 at the age of 38.

We wish to suggest a structure for the salt of deoxyribose nucleic acid (DNA). . . . The structure has two helical chains each coiled round the same axis. . . . Both chains follow right-handed helices, but the two chains run in opposite directions. . . . The bases are on the inside of the helix and the phosphates on the outside. . . . There is a residue on each chain every 3.4 Å and the structure repeats after 10 residues. . . . The novel feature of the structure is the manner in which the two chains are held together by the purine and pyrimidine bases. The planes of the bases are perpendicular to the fiber axis. They are joined together in pairs, a single base from one chain being hydrogen-bonded to a single base from the other chain, so that the two lie side by side. One of the pair must be a purine and the other a pyrimidine for bonding to occur. . . . Only specific pairs of bases can bond together. These pairs are adenine (purine) with thymine (pyrimidine), and guanine (purine) with cytosine (pyrimidine). In other words, if an adenine forms one member of a pair, on either chain, then on these assumptions the other member must be thymine; similarly for gua-

nine and cytosine. The sequence of bases on a single chain does not appear to be restricted in any way. However, if only specific pairs of bases can be formed, it follows that if the sequence of bases on one chain is given, then the se-

quence on the other chain is automatically determined. . . . It has not escaped our notice that the specific pairing we have postulated immediately suggests a plausible copying mechanism for the genetic material. . . . We are much indebted to Dr. Jerry Donohue for constant advice and criticism, especially on

If only specific pairs of bases can be formed, it follows that if the sequence of bases on one chain is given, then the sequence on the other chain is automatically determined.

interatomic distances. We have also been stimulated by a knowledge of the general nature of the unpublished experimental results and ideas of Dr. Maurice H. F. Wilkins, Dr. Rosalind Franklin and their co-workers at King's College, London.

Source: *Nature* 171: 737–738

DNA Structure as Related to Function

In the structure of the DNA molecule, we can see how three essential requirements of a genetic material are met.

1. Any genetic material must be able to be replicated accurately, so that the information it contains will be precisely replicated and inherited by daughter cells. The basis for exact duplication of a DNA molecule is the

pairing of A with T and of G with C in the two polynucleotide chains. Unwinding and separation of the chains, with each free chain being copied, results in the formation of two identical double helices (see Figure 1.6).

2. A genetic material must also have the capacity to carry all of the information needed to direct the organization and metabolic activities of the cell. As we saw in Chapter 1, the product of most genes is a protein molecule—a polymer composed of repeating units of amino acids. The sequence

of amino acids in the protein determines its chemical and physical properties. A gene is expressed when its protein product is synthesized, and one requirement of the genetic material is that it direct the order in which amino acid units are added to the end of a growing protein molecule. In DNA, this is done by means of a genetic code in which groups of three bases specify amino acids. Because the four bases in a DNA molecule can be arranged in any sequence, and because the sequence can vary from one part of the molecule to another and from organism to organism, DNA can contain a great many unique regions, each of which can be a distinct gene. A long DNA chain can direct the synthesis of a variety of different protein molecules.

3. A genetic material must also be capable of undergoing occasional mutations in which the information it carries is altered. Furthermore, so that mutations will be heritable, the mutant molecules must be capable of being replicated as faithfully as the parental molecule. This feature is necessary to account for the evolution of diverse organisms through the slow accumulation of favorable mutations. Watson and Crick suggested that heritable mutations might be possible in DNA by rare mispairing of the bases, with the result that an incorrect nucleotide becomes incorporated into a replicating DNA strand.

2.3 The Separation and Identification of Genomic DNA Fragments

The following sections show how an understanding of DNA structure and replication has been put to practical use in the development of procedures for the separation and identification of particular DNA fragments. These methods are used primarily either to identify DNA markers or to aid in the isolation of particular DNA fragments that are of genetic interest. For example, consider a pedigree of familial breast cancer in which a particular DNA fragment serves as a marker for a bit of chromosome that also includes the mutant gene responsible for the increased risk; then the ability to identify the fragment is critically important in assessing the relative risk for each of the women in the pedigree

who might carry the mutant gene. To take another example, suppose there is reason to believe that a mutation causing a genetic disease is present in a particular DNA fragment; then it is important to be able to pinpoint this fragment and isolate it from affected individuals to verify whether this hypothesis is true and, if so, to identify the nature of the mutation.

Most procedures for the separation and identification of DNA fragments can be grouped into two general categories:

1. Those that identify a specific DNA fragment present in genomic DNA by making use of the fact that complementary single-stranded DNA sequences can, under the proper conditions, form a duplex molecule. These procedures rely on *nucleic acid hybridization*.
2. Those that use prior knowledge of the sequence at the ends of a DNA fragment to specifically and repeatedly replicate this one fragment from genomic DNA. These procedures rely on selective DNA replication (*amplification*) by means of the *polymerase chain reaction*.

The major difference between these approaches is that the first (relying on nucleic acid hybridization) identifies fragments that are present in the genomic DNA itself, whereas the second (relying on DNA amplification) identifies experimentally manufactured *replicas* of fragments whose original templates (but not the replicas) were present in the genomic DNA. This difference has practical implications:

- Hybridization methods require a greater amount of genomic DNA for the experimental procedures, but relatively large fragments can be identified, and no prior knowledge of the DNA sequence is necessary.
- Amplification methods require extremely small amounts of genomic DNA for the experimental procedures, but the amplification is usually restricted to relatively small fragments, and some prior knowledge of DNA sequence is necessary.

The following sections discuss both types of approaches and give examples of how they are used. In methods that use nucleic acid hybridization to identify particular fragments present in genomic DNA, the first step is usually cutting the genomic DNA into fragments of experimentally manageable size. This procedure is discussed next.

Restriction Enzymes and Site-Specific DNA Cleavage

Procedures for chemical isolation of DNA, such as those developed by Avery, MacLeod, and McCarty (Chapter 1), usually lead to random breakage of double-stranded molecules into an average length of about 50,000 base pairs. This length is denoted 50 kb, where kb stands for **kilobases** (1 kb = 1000 base pairs). A length of 50 kb is close to the length of double-stranded DNA present in the bacteriophage λ that infects *E. coli*. The 50-kb fragments can be made shorter by vigorous shearing forces, such as occur in a kitchen blender, but one of the problems with breaking large DNA molecules into smaller fragments by random shearing is that the fragments containing a particular gene, or part of a gene, will be of different sizes. In other words, with random shearing, it is not possible to isolate and identify a *particular* DNA fragment on the basis of its size and sequence content, because each randomly sheared molecule that contains the desired sequence somewhere within it differs in size from all other molecules that contain the sequence. In this section we describe an important enzymatic technique that can be used for cleaving DNA molecules at specific sites. This method ensures that all DNA fragments that contain a particular sequence have the same size; furthermore, each fragment that contains the desired sequence has the sequence located at exactly the same position within the fragment.

The cleavage method makes use of an important class of DNA-cleaving enzymes isolated primarily from bacteria. The enzymes are called **restriction endonucleases** or **restriction enzymes**, and they are able to cleave DNA molecules at the positions at which particular, short sequences of bases are present. These naturally occurring enzymes serve to protect the bacterial cell by disabling the DNA of bacteriophages that attack it. Their discovery earned Werner Arber of Switzerland a Nobel Prize in 1978. Technically, the enzymes are known as *type II restriction endonucleases*. The restriction enzyme *Bam*HI is one example; it recognizes the double-stranded sequence

5'-GGATCC-3'
3'-CCTAGG-5'



Figure 2.9 Structure of the part of the restriction enzyme *Bam*HI that comes into contact with its recognition site in the DNA (blue). The pink and green cylinders represent regions of the enzyme in which the amino acid chain is twisted in the form of a right-handed helix. [Courtesy of A. A. Aggarwal. Reprinted with permission from M. Newman, T. Strzelecka, L. F. Dorner, I. Schildkraut, and A. A. Aggarwal, 1995. *Science* 269: 656. Copyright 2000 American Association for the Advancement of Science.]

and cleaves each strand between the G-bearing nucleotides shown in red. Figure 2.9 shows how the regions that make up the active site of *Bam*HI contact the recognition site (blue) just prior to cleavage, and the cleavage reaction is indicated in Figure 2.10.

Table 2.3 lists nine of the several hundred restriction enzymes that are known. Most restriction enzymes are named after the species in which they were found. *Bam*HI, for example, was isolated from the bacterium *Bacillus amyloliquefaciens* strain H, and it is the first (I) restriction enzyme isolated from this organism. Because the first three letters in the name of each restriction enzyme stand for the bacterial species of origin, these letters are printed in *italics*; the rest of the symbols in the name are not italicized. Most restriction enzymes recognize only one short base sequence, usually four or six nucleotide pairs. The enzyme binds with the DNA at these sites and makes a break in each strand of the DNA molecule, producing 3'-OH and 5'-P groups at each position. The nucleotide sequence recognized for cleavage by a restriction enzyme is called the **restriction site** of the enzyme. The examples in Table 2.3 show that some restriction enzymes cleave their restriction site asymmetrically (at different sites in the

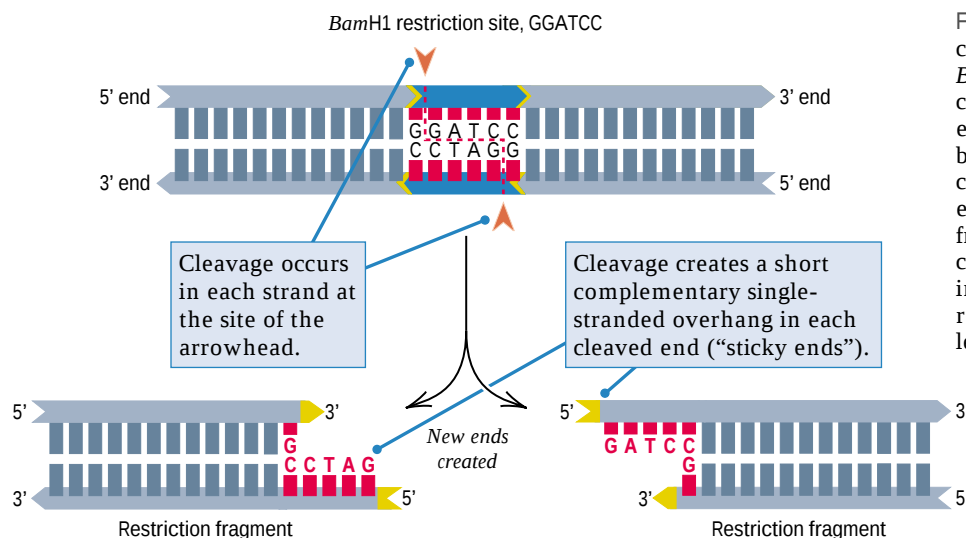
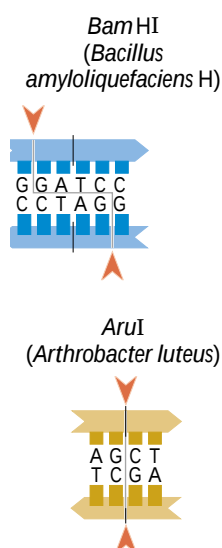


Figure 2.10 Mechanism of DNA cleavage by the restriction enzyme *Bam*HI. Wherever the duplex contains a *Bam*HI restriction site, the enzyme makes a single cut in the backbone of each DNA strand. Each cut creates a new 3' end and a new 5' end, separating the duplex into two fragments. In the case of *Bam*HI the cuts are staggered cuts, so the resulting ends terminate in single-stranded regions, each four base pairs in length.

Table 2.3 Some restriction endonucleases, their sources, and their cleavage sites

Enzyme (Microorganism)	Enzyme (Microorganism)	Enzyme (Microorganism)
<i>Eco</i> RI (<i>Escherichia coli</i>)	<i>Hind</i> III (<i>Haemophilus influenzae</i>)	<i>Aru</i> I (<i>Arthrobacter luteus</i>)
<i>Bam</i> HI (<i>Bacillus amyloliquefaciens</i> H)	<i>Pst</i> I (<i>Providencia stuartii</i>)	<i>Rsa</i> I (<i>Rhodopseudomonas sphaeroides</i>)
<i>Hae</i> II (<i>Haemophilus aegyptus</i>)	<i>Taq</i> I (<i>Thermus aquaticus</i>)	<i>Pvu</i> II (<i>Proteus vulgaris</i>)

Note: The vertical dashed line indicates the axis of symmetry in each sequence. Red arrows indicate the sites of cutting. The enzyme *Taq*I yields cohesive ends consisting of two nucleotides, whereas the cohesive ends produced by the other enzymes contain four nucleotides. Pu and Py refer to any purine and pyrimidine, respectively.



two DNA strands), but other restriction enzymes cleave symmetrically (at the same site in both strands). The former leave **sticky ends** because each end of the cleaved site has a small, single-stranded overhang that is complementary in base sequence to the other end (Figure 2.10). In contrast, enzymes that have symmetrical cleavage sites yield DNA fragments that have **blunt ends**. In virtually all cases, the restriction site of a restriction enzyme reads the same on both strands, provided that the opposite polarity of the strands is taken into account; for example, each strand in the restriction site of *Bam*HI reads 5'-GGATCC-3' (Figure 2.10). A DNA sequence with this type of symmetry is called a **palindrome**. (In ordinary English, a palindrome is a word or phrase that reads the same forwards and backwards, such as “madam.”)

Restriction enzymes have the following important characteristics:

- Most restriction enzymes recognize a single restriction site.
- The restriction site is recognized without regard to the source of the DNA.
- Because most restriction enzymes recognize a unique restriction site sequence, the number of cuts in the DNA from a particular organism is determined by the number of restriction sites present.

The DNA fragment produced by a pair of adjacent cuts in a DNA molecule is called a **restriction fragment**. A large DNA molecule will typically be cut into many restriction fragments of different sizes. For example, an *E. coli* DNA molecule, which contains 4.6×10^6 base pairs, is cut into several hundred to several thousand fragments, and mammalian genomic DNA is cut into more than a million fragments. Although these numbers are large, they are actually quite small relative to the number of sugar–phosphate bonds in the DNA of an organism.

Gel Electrophoresis

The DNA fragments produced by a restriction enzyme can be separated by size using the fact that DNA is negatively charged and moves in response to an electric field. If the terminals of an electrical power source are connected to the opposite ends of a horizontal tube containing a DNA solution, then the DNA molecules will move toward the positive end of the tube at a rate that depends on the electric field strength and on the shape and size of the molecules. The movement of charged molecules in an electric field is called *electrophoresis*.

The type of electrophoresis most commonly used in genetics is **gel electrophoresis**. An experimental arrangement for gel electrophoresis of DNA is shown in Figure 2.11. A thin slab of a gel, usually agarose or acrylamide, is prepared containing small slots (called *wells*) into which samples are placed. An electric field is applied, and the negatively charged DNA molecules penetrate and move through the gel toward the anode (the positively charged electrode). A gel is a complex molecular network that contains narrow, tortuous passages, so smaller DNA molecules pass through more easily; hence the rate of movement increases as the size of the DNA fragment decreases. Figure 2.12 shows the result of electrophoresis of a set of double-stranded DNA molecules in an agarose gel. Each discrete region containing DNA is called a **band**. The bands can be visualized under ultraviolet light after soaking the gel in the dye *ethidium bromide*, the molecules of which intercalate between the stacked bases in duplex DNA and render it fluorescent. In Figure 2.12, each band in the gel

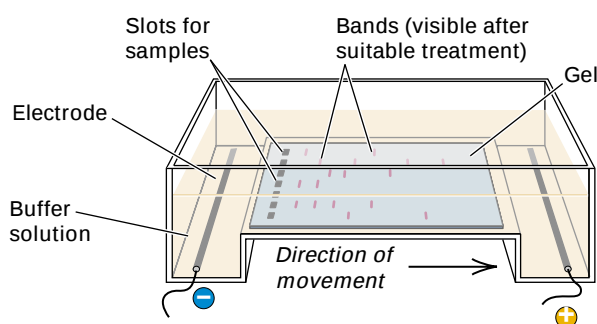


Figure 2.11 Apparatus for gel electrophoresis. Liquid gel is allowed to harden with an appropriately shaped mold in place to form “wells” for the samples (purple). After electrophoresis, the DNA fragments, located at various positions in the gel, are made visible by immersing the gel in a solution containing a reagent that binds to or reacts with DNA. The separated fragments in a sample appear as bands, which may be either visibly colored or fluorescent, depending on the particular reagent used. The region of a gel in which the fragments in one sample can move is called a *lane*; this gel has seven lanes.

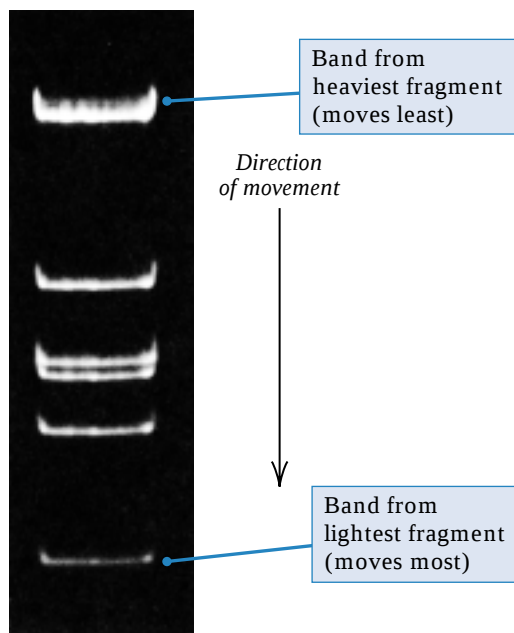


Figure 2.12 Gel electrophoresis of DNA. Fragments of different sizes were mixed and placed in a well. Electrophoresis was in the downward direction. The DNA has been made visible by the addition of a dye (ethidium bromide) that binds only to DNA and that fluoresces when the gel is illuminated with short-wavelength ultraviolet light.

results from the fact that all DNA fragments of a given size have migrated to the same position in the gel. To produce a visible band, a minimum of about 5×10^{-9} grams of DNA is required, which for a fragment of size 3 kb works out to about 10^9 molecules. The point is that a very large number of copies of any particular DNA fragment must be present in order to yield a visible band in an electrophoresis gel.

A linear double-stranded DNA fragment has an electrophoretic mobility that decreases in proportion to the logarithm of its length in base pairs—the longer the fragment, the slower it moves—but the proportionality constant depends on the agarose concentration, the composition of the buffering solution, and the electrophoretic conditions. This means that different concentrations of agarose allow efficient separation of different size ranges of DNA fragments (see Figure 2.13). Less dense gels, such as 0.6 percent agarose, are used to separate larger fragments; whereas more dense gels, such as 2 percent agarose, are used to separate smaller fragments. The inset in Figure 2.13 shows the dependence of electrophoretic mobility on the logarithm of fragment size. It also indicates that the linear relationship breaks down for the largest fragments that can be resolved under a given set of conditions.

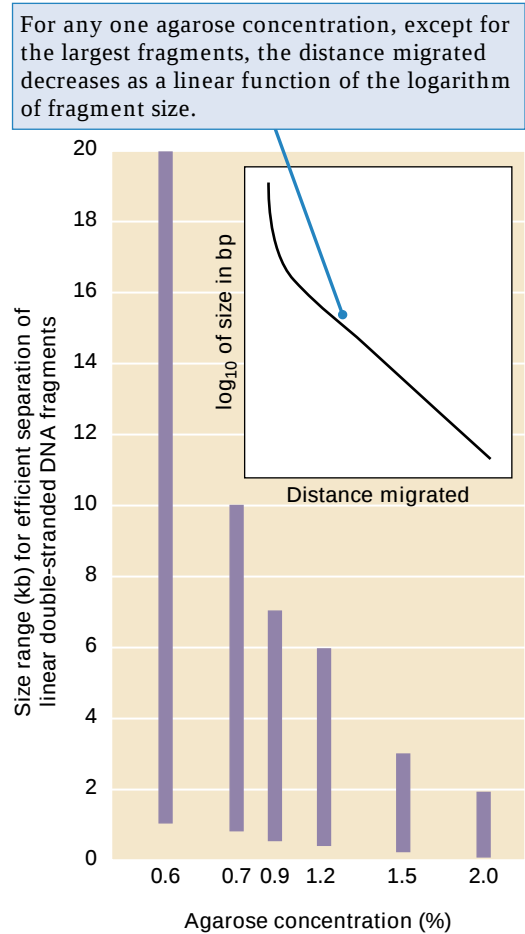


Figure 2.13 In agarose gels, the concentration of agarose is an important factor in determining the size range of DNA fragments that can be separated.

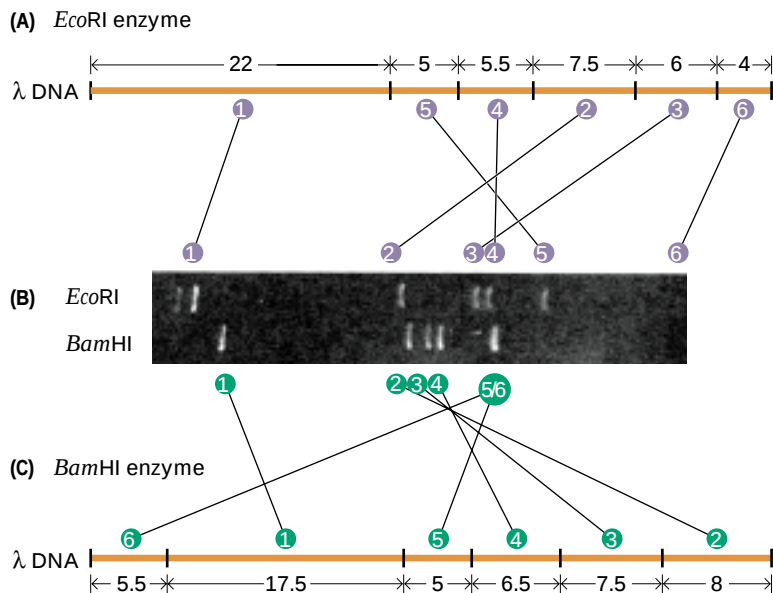


Figure 2.14 Restriction maps of λ DNA for the restriction enzymes (A) *EcoRI* and (C) *BamHI*. The vertical bars indicate the sites of cutting. The numbers within the arrows are the approximate lengths of the fragments in kilobase pairs (kb). (B) An electrophoresis gel of *BamHI* and *EcoRI* enzyme digests of λ DNA. Numbers indicate fragments in order from largest (1) to smallest (6); the circled numbers on the maps correspond to the numbers beside the gel. The DNA has not undergone electrophoresis long enough to separate bands 5 and 6 of the *BamHI* digest. *Note:* In Problem 2 at the end of this chapter (Guide to Problem Solving), we show how to use the results of a double digest to determine the particular order of fragments for a pair of restriction enzymes.

Because of the sequence specificity of cleavage, a particular restriction enzyme produces a unique set of fragments for a particular DNA molecule. Another enzyme will produce a different set of fragments from the same DNA molecule. In Figure 2.14, this principle is illustrated for the digestion of *E. coli* phage λ DNA by either *EcoRI* or *BamHI* (see part B). The locations of the cleavage sites for these enzymes in λ DNA are shown in Figures 2.14A and C. A diagram showing sites of cleavage along a DNA molecule is called a **restriction map**. Particular DNA fragments can be isolated by cutting out the small region of the gel that contains the fragment and removing the DNA from the gel. One important use of isolated restriction fragments employs the enzyme *DNA ligase* to insert them into self-replicating molecules such as bacteriophage, plasmids, or even small artificial chromosomes (Figure 2.2). These procedures constitute **DNA cloning** and are the basis of one form of

genetic engineering, discussed further in Chapter 13.

DNA fragments that have been cloned into organisms such as *E. coli* are widely used because the fragments can be isolated in large amounts and purified relatively easily. Among the uses of cloned DNA are:

- *DNA sequencing.* All current methods of DNA sequencing require cloned DNA fragments. These methods are discussed in Chapter 6.
- *Nucleic acid hybridization.* As we shall see below, an important application of cloned DNA fragments entails incorporating a radioactive or light-emitting “label” into them, after which the labeled material is used to “tag” DNA fragments containing similar sequences.
- *Storage and distribution.* Cloned DNA can be stored for long periods without risk of change and can easily be distributed to other researchers.

Nucleic Acid Hybridization

Most genomes are sufficiently large and complex that digestion with a restriction enzyme produces many bands that are the same or similar in size. Identifying a particular DNA fragment in a background of many other fragments of similar size presents a needle-in-a-haystack problem. Suppose, for example, that we are interested in a particular 3.0 *BamHI* fragment from the human genome that serves as a marker indicating the presence of a genetic risk factor toward breast cancer among the individuals in a particular pedigree. This fragment of 3.0 kb is indistinguishable, on the basis of size alone, from fragments ranging from about 2.9 to 3.1 kb. How many fragments in this size range are expected? When human genomic DNA is cleaved with *BamHI*, the average length of a restriction fragment is $4^6 = 4096$ base pairs, and the expected total number of *BamHI* fragments is about 730,000; in the size range 2.9–3.1 kb, the expected number of fragments is about 17,000. What this means is that even though we know that the fragment we are interested in is 3 kb in length, it is only one of 17,000 fragments that are so similar in size that ours cannot be distinguished from the others by length alone.

This identification task is actually harder than finding a needle in a haystack because

haystacks are usually dry. A more accurate analogy would be looking for a needle in a haystack that had been pitched into a swimming pool full of water. This analogy is more relevant because gels, even though they contain a supporting matrix to make them semisolid, are primarily composed of water, and each DNA molecule within a gel is surrounded entirely by water. Clearly, we need some method by which the molecules in a gel can be immobilized and our *specific* fragment identified.

The DNA fragments in a gel are usually immobilized by transferring them onto a sheet of special filter paper consisting of nitrocellulose, to which DNA can be permanently (covalently) bound. How this is done is described in the next section. In this section we examine how the two strands in a double helix can be “unzipped” to form single strands and how, under the proper conditions, two single strands that are complementary or nearly complementary in sequence can be “zipped” together to form a different double helix. The “unzipping” is called **denaturation**, the “zipping” is called **renaturation**. The practical applications of denaturation and renaturation are many:

- A small part of a DNA fragment can be “zipped” with a much larger DNA fragment. This principle is used in identifying specific DNA fragments in a complex mixture, such as the 3-kb *Bam*HI marker for breast cancer that we have been considering. Applications of this type include the tracking of genetic markers in pedigrees and the isolation of fragments containing a particular mutant gene.
- A DNA fragment from one gene can be “zipped” with similar fragments from other genes in the same genome; this principle is used to identify different members of *families* of genes that are similar, but not identical, in sequence and that have related functions.
- A DNA fragment from one species can be “zipped” with similar sequences from other species. This allows the isolation of genes that have the same or related functions in multiple species. It is used to study aspects of molecular evolution, such as how differences in sequence are correlated with differences in function, and the patterns and rates of change in gene sequences as they evolve.

As we saw in Section 2.2, the double-stranded helical structure of DNA is maintained by base stacking and by hydrogen bonding between the complementary base

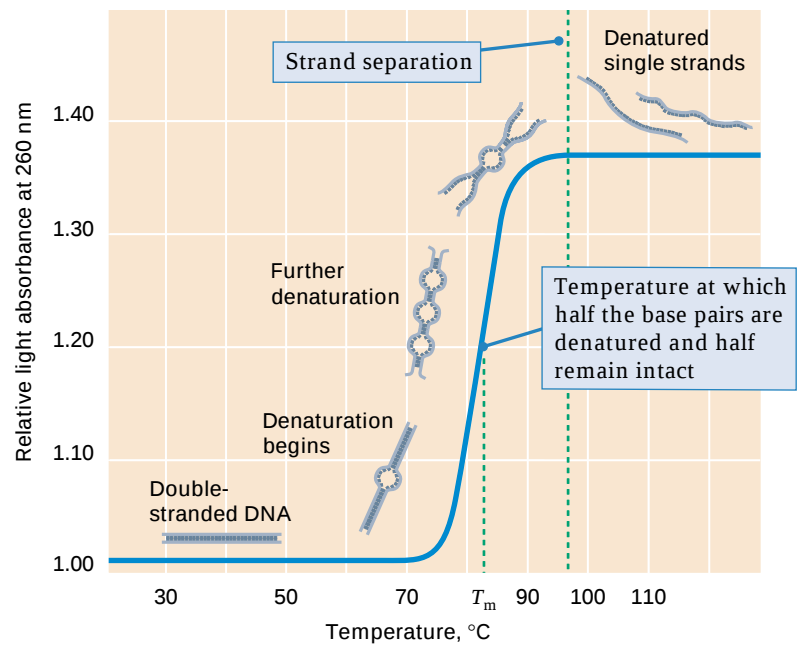


Figure 2.15 Mechanism of denaturation of DNA by heat. The temperature at which 50 percent of the base pairs are denatured is the *melting temperature*, symbolized T_m .

pairs. When solutions containing DNA fragments are raised to temperatures in the range 85–100°C, or to the high pH of strong alkaline solutions, the paired strands begin to separate, or “unzip.” Unwinding of the helix happens in less than a few minutes (the time depends on the length of the molecule). When the helical structure of DNA is disrupted, and the strands have become completely unzipped, the molecule is said to be **denatured**. A common way to detect denaturation is by measuring the capacity of DNA in solution to absorb ultraviolet light of wavelength 260 nm, because the absorption at 260 nm (A_{260}) of a solution of single-stranded molecules is 37 percent higher than the absorption of the double-stranded molecules at the same concentration. As shown in Figure 2.15, the progress of denaturation can be followed by slowly heating a solution of double-stranded DNA and recording the value of A_{260} at various temperatures. The temperature required for denaturation increases with $G + C$ content, not only because $G-C$ base pairs have three hydrogen bonds and $A-T$ base pairs two, but because consecutive $G-C$ base pairs have stronger base stacking.

09131_01_1746P

Denatured DNA strands can, under certain conditions, form double-stranded DNA with other strands, provided that the strands are sufficiently complementary in sequence. This process of renaturation is called **nucleic acid hybridization** because the double-stranded molecules are “hybrid” in that each strand comes from a different source. For DNA strands to hybridize, two requirements must be met:

1. The salt concentration must be high ($> 0.25\text{M}$) to neutralize the negative charges of the phosphate groups, which would otherwise cause the complementary strands to repel one another.
2. The temperature must be high enough to disrupt hydrogen bonds that form at random between short sequences of bases within the same strand, but not so high that stable base pairs between the complementary strands are disrupted.

The initial phase of renaturation is a slow process because the rate is limited by the random chance that a region of two complementary strands will come together to form a short sequence of correct base pairs. This initial pairing step is followed by a rapid pairing of the remaining complementary bases and rewinding of the helix. Rewinding is accomplished in a matter of seconds, and its rate is independent of DNA concentration because the complementary strands have already found each other.

The example of nucleic acid hybridization in Figure 2.16 will enable us to understand some of the molecular details and also to see how hybridization is used to “tag” and identify a particular DNA fragment.

Shown in part A is a solution of denatured DNA, called the **probe**, in which each molecule has been labeled with either radioactive atoms or light-emitting molecules. Probe DNA is typically obtained from a clone, and the labeled probe usually contains denatured forms of both strands present in the original duplex molecule. (This has led to some confusing terminology. Geneticists say that probe DNA hybridizes with DNA fragments containing sequences that are *similar* to the probe, rather than *complementary*. What actually occurs is that one strand of the probe undergoes hybridization with a complementary sequence in the fragment. But because the probe usually contains both strands, hybridization takes place with any fragment that contains a similar sequence, each strand in the probe undergoing hybridization with the complementary sequence in the fragment.)

Part B in Figure 2.16 is a diagram of genomic DNA fragments that have been immobilized on a nitrocellulose filter. When the probe is mixed with the genomic fragments (part C), random collisions bring short, complementary stretches together. If the region of complementary sequence is short (part D), then random collision cannot initiate renaturation because the flanking sequences cannot pair; in this case the probe falls off almost immediately. If, however, a collision brings short sequences together in the correct register (part E), then this initiates renaturation, because the pairing proceeds zipperlike from the initial contact. The main point is that DNA fragments are able to hybridize only if the length of

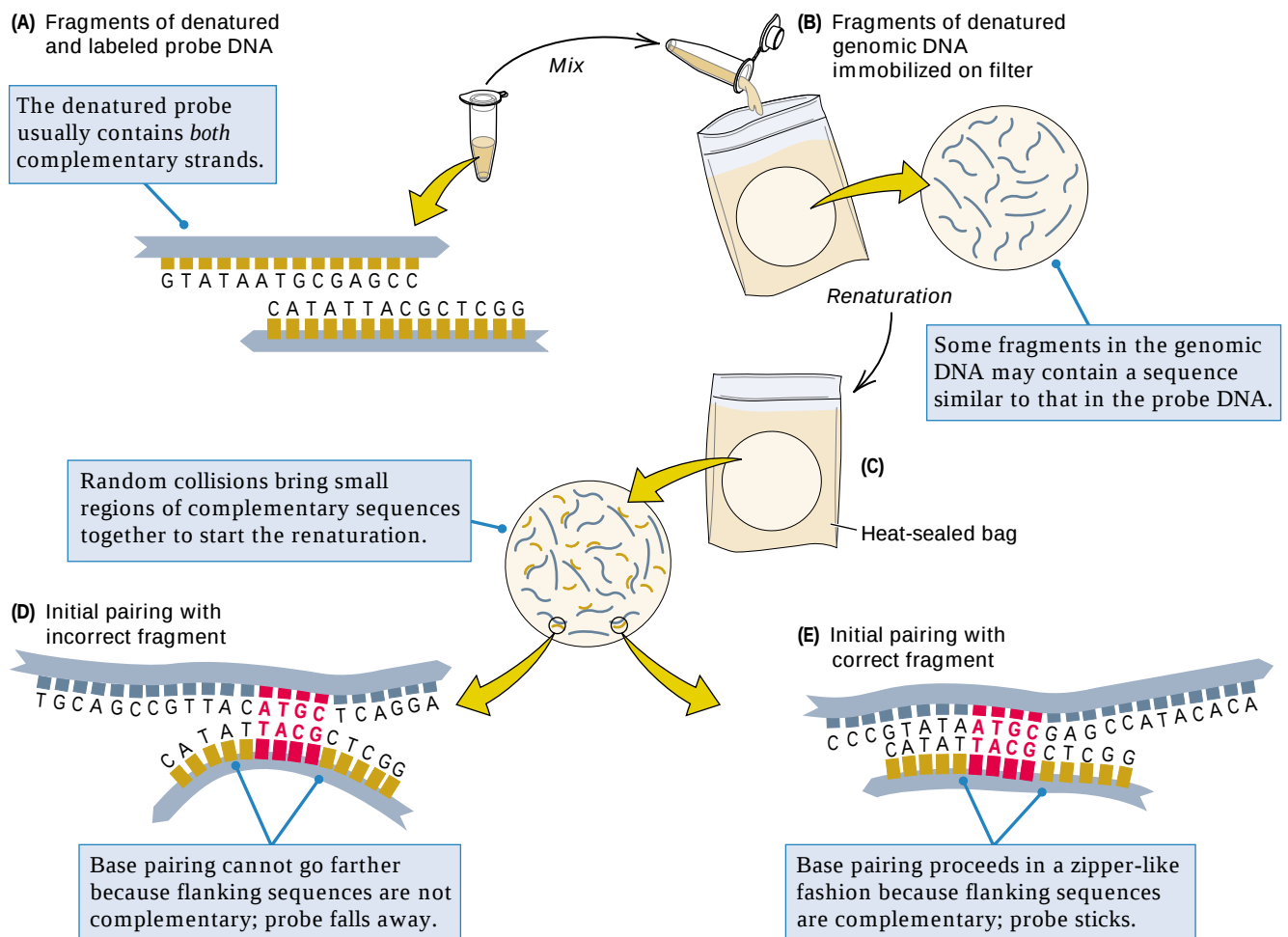


Figure 2.16 Nucleic acid hybridization. (A) Duplex molecules of probe DNA (obtained from a clone) are denatured and (B) placed in contact with a filter to which is attached denatured strands of genomic DNA. (C) Under the proper conditions of salt concentration and temperature, short complementary stretches come together by random collision. (D) If the sequences flanking the paired region are not complementary, then the pairing is unstable and the strands come apart again. (E) If the sequences flanking the paired region are complementary, then further base pairing stabilizes the renatured duplex.

the region in which they can pair is sufficiently long. Some mismatches in the paired region can be tolerated. How many mismatches are allowed is determined by the conditions of the experiment: The lower the temperature at which the hybridization is carried out, and the higher the salt concentration, the greater the proportion of mismatches that are tolerated.

The Southern Blot

The ability to renature DNA in the manner outlined in Figure 12.16 means that solution containing a small fragment of dena-

tured DNA, if it is suitably labeled (for example, with radioactive ^{32}P), can be combined with a complex mixture of denatured DNA fragments, and upon renaturation the small fragment will “tag” with radioactivity any molecules in the complex mixture with which it can hybridize. The radioactive tag allows these molecules to be identified.

The methods of DNA cleavage, electrophoresis, transfer to nitrocellulose, and hybridization with a probe are all combined in the **Southern blot**, named after its inventor Edward Southern. In this procedure, a gel in which DNA molecules have been separated by electrophoresis is treated with

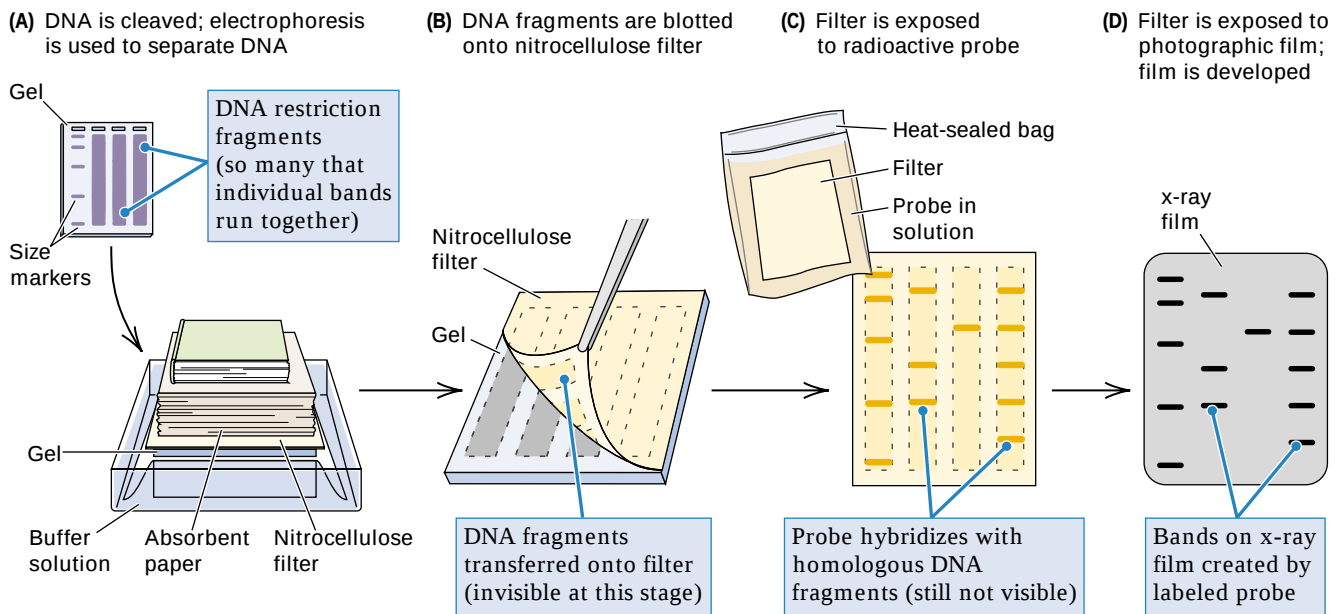


Figure 2.17 Southern blot. (A) DNA restriction fragments are separated by electrophoresis, blotted from the gel onto a nitrocellulose or nylon filter, and chemically attached by the use of ultraviolet light. (B) The strands are denatured and mixed with radioactive or light-sensitive probe DNA, which binds with complementary sequences present on the filter. The bound probe remains, whereas unbound probe washes off. (C) Bound probe is revealed by darkening of photographic film placed over the filter. The positions of the bands indicate which restriction fragments contain DNA sequences homologous to those in the probe.

alkali to denature the DNA and render it single-stranded (Figure 2.17). Then the DNA is transferred to a sheet of nitrocellulose filter in such a way that the relative positions of the DNA fragments are maintained. The transfer is accomplished by overlaying the nitrocellulose onto the gel and stacking many layers of absorbent paper on top; the absorbent paper sucks water molecules from the gel and through the nitrocellulose, to which the DNA fragments adhere. (This step is the “blot” component of the Southern blot, parts A and B.) Then the filter is treated so that the single-stranded DNA becomes permanently bound. The treated filter is mixed with a solution containing denatured probe (DNA or RNA) under conditions that allow complementary strands to hybridize to form duplex molecules (part C). Radioactive or other label present in the probe becomes stably bound to the filter, and therefore resistant to removal by washing, only at positions at which base sequences complementary to the probe are already present on the filter, so that the probe can form duplex molecules. The label is located by placing the paper in contact with x-ray film. After development of the film, blackened regions indicate positions of bands containing the radioactive or light-emitting label (part D).

The procedure in Figure 2.17 solves the wet-haystack problem by transferring and immobilizing the genomic DNA fragments to a filter and identifying, by hybridization,

the ones that are of interest. Practical applications of Southern blotting center on identifying DNA fragments that contain sequences similar to the probe DNA or RNA, where the proportion of mismatched nucleotides allowed is determined by the conditions of hybridization. The advantages of the Southern blot are convenience and sensitivity. The sensitivity comes from the fact that both hybridization with a labeled probe and the use of photographic film amplify the signal; under typical conditions, a band can be observed on the film with only 5×10^{-12} grams of DNA—a thousand times less DNA than the amount required to produce a visible band in the gel itself.

2.4 Selective Replication of Genomic DNA Fragments

Although nucleic acid hybridization allows a particular DNA fragment to be identified when present in a complex mixture of fragments, it does not enable the fragment to be separated from the others and purified. Obtaining the fragment in purified form requires cloning, which is straightforward but time-consuming. (Cloning methods are discussed in Chapter 13.) However, if the fragment of interest is not too long, and if the nucleotide sequence at each end is known, then it becomes possible to obtain large quantities of the fragment merely by

Au: Proofreader asks that you confirm reference to Chapter 13 in first para of sec. 2.4.

selective replication. This process is called **amplification**. How would one know the sequence of the ends? Let us return to our example of the 3.0-kb *Bam*HI fragment that serves to mark a risk factor for breast cancer in certain pedigrees. Suppose that this fragment is cloned and sequenced from one affected individual, and it is found that, relative to the normal genomic sequence in this region, the *Bam*HI fragment is missing a region of 500 base pairs. At this point the sequences at the ends of the fragment are known, and we can also infer that amplification of genomic DNA from individuals with the risk factor will yield a band of 3.0 kb, whereas amplification from genomic DNA of noncarriers will yield a band of 3.5 kb. This difference allows every person in the pedigree to be diagnosed as a carrier or noncarrier merely by means of DNA amplification. To understand how amplification works, it is first necessary to examine a few key features of DNA replication.

Constraints on DNA Replication: Primers and 5'-to-3' Strand Elongation

As with most metabolic reactions in living cells, nucleic acids are synthesized in chemical reactions controlled by enzymes. An enzyme that forms the sugar-phosphate bond (the phosphodiester bond) between adjacent nucleotides in a nucleic acid chain is called a **DNA polymerase**. A variety of DNA polymerases have been purified, and for amplification of a DNA fragment, the DNA synthesis is carried out *in vitro* by combining purified cellular components in a test tube under precisely defined conditions. (*In vitro* means “without participation of living cells.”)

In order for DNA polymerase to catalyze synthesis of a new DNA strand, preexisting single-stranded DNA must be present. Each single-stranded DNA molecule present in the reaction mix can serve as a template upon which a new partner strand is created by the DNA polymerase. For DNA replication to take place, the 5'-triphosphates of the four deoxynucleosides must also be present. This requirement is rather obvious, because the nucleoside triphosphates are the precursors from which new DNA strands are created. The triphosphates needed are the compounds denoted in Table 2.1 as

dATP, dGTP, dTTP, and dCTP, which contain the bases adenine, guanine, thymine, and cytosine, respectively. Details of the structures of dCTP and dGTP are shown in Figure 2.18, in which the phosphate groups cleaved off during DNA synthesis are indicated. DNA synthesis requires all four nucleoside 5'-triphosphates and does not take place if any of them is omitted.

A feature found in all DNA polymerases is that

A DNA polymerase can only *elongate* a DNA strand. It is not possible for DNA polymerase to *initiate* synthesis of a new strand, even when a template molecule is present.

One important implication of this principle is that DNA synthesis requires a pre-existing segment of nucleic acid that is hydrogen-bonded to the template strand. This segment is called a **primer**. Because the primer molecule can be very short, it is an **oligonucleotide**, which literally means “few nucleotides.” As we shall see in Chapter 6, in living cells the primer is a short segment of RNA, but in DNA amplification *in vitro*, the primer employed is usually DNA.

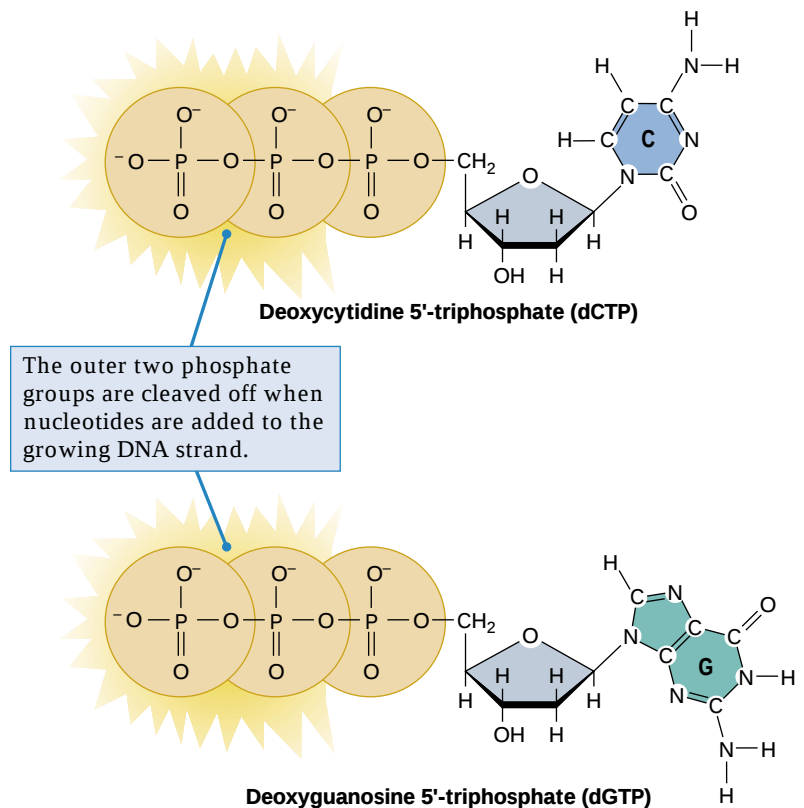


Figure 2.18 Two deoxynucleoside triphosphates used in DNA synthesis. The outer two phosphate groups are removed during synthesis.

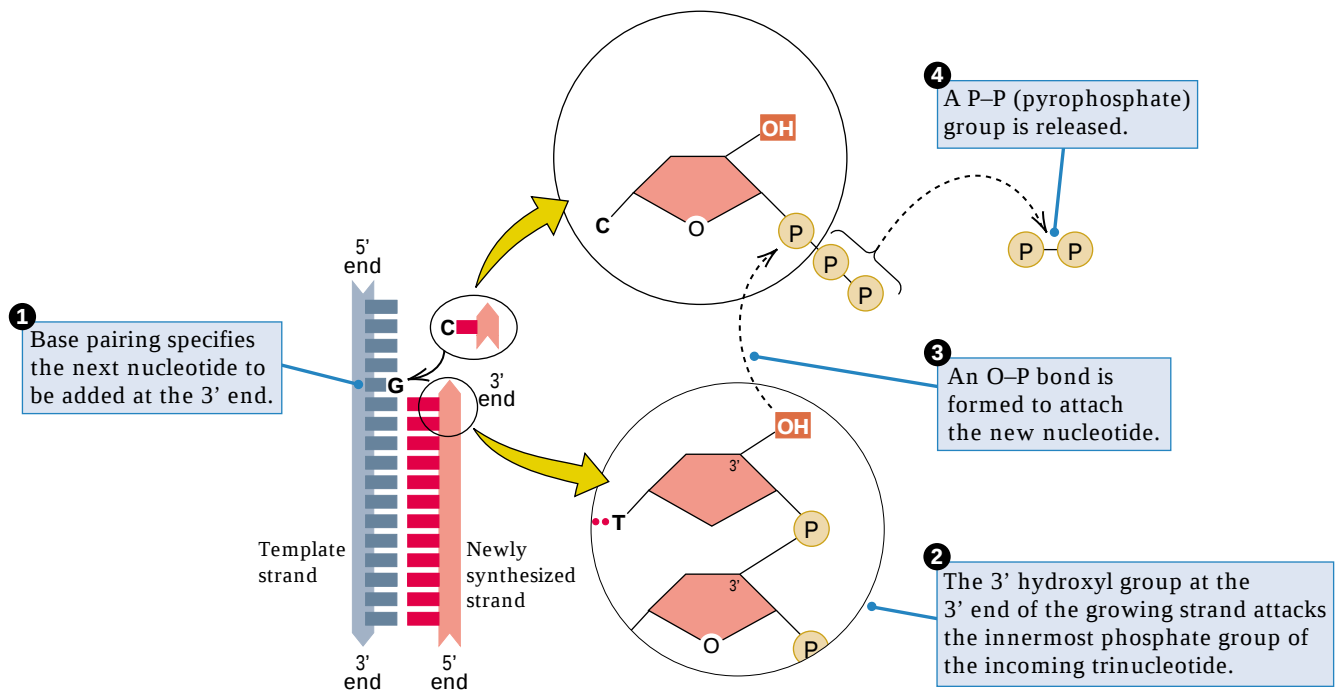


Figure 2.19 Addition of nucleotides to the 3'-OH terminus of a growing strand. The recognition step is shown as the formation of hydrogen bonds between the A and the T. The chemical reaction is that the 3'-OH group of the 3' end of the growing chain attacks the innermost phosphate group of the incoming trinucleotide.

It is the 3' end of the primer that is essential, because, as emphasized in Chapter 1,

DNA synthesis proceeds only by addition of successive nucleotides to the 3' end of the growing strand. In other words, chain elongation always takes place in the 5'-to-3' direction (5' → 3').

The reason for the 5' → 3' direction of chain elongation is illustrated in Figure 2.19. It is a consequence of the fact that the reaction catalyzed by DNA polymerase is the formation of a phosphodiester bond between the free 3'-OH group of the chain being extended and the innermost phosphorus atom of the nucleoside triphosphate being incorporated at the 3' end. Recognition of the appropriate incoming nucleoside triphosphate in replication depends on base pairing with the opposite nucleotide in the template strand. DNA polymerase will usually catalyze the polymerization reaction that incorporates the new nucleotide at the primer terminus only when the correct base pair is present. The same DNA polymerase is used to add each of the four deoxynucleoside phosphates to the 3'-OH terminus of the growing strand.

The Polymerase Chain Reaction

The requirement for an oligonucleotide primer, and the constraint that chain elongation must always occur in the 5' → 3' direction, make it possible to obtain large quantities of a particular DNA sequence by selective amplification *in vitro*. The method for selective amplification is called the **polymerase chain reaction (PCR)**. For its invention, Californian Kary B. Mullis was awarded a Nobel Prize in 1993. PCR amplification uses DNA polymerase and a *pair* of short, synthetic oligonucleotide primers, usually 18–22 nucleotides in length, that are complementary in sequence to the ends of the DNA sequence to be amplified. Figure 2.20 gives an example in which the primer oligonucleotides (green) are 9-mers. These are too short for most practical purposes, but they will serve for illustration. The original duplex molecule (part A) is shown in blue. This duplex is mixed with a vast excess of primer molecules, DNA polymerase, and all four nucleoside triphosphates. When the temperature is raised, the strands of the duplex denature and become separated.

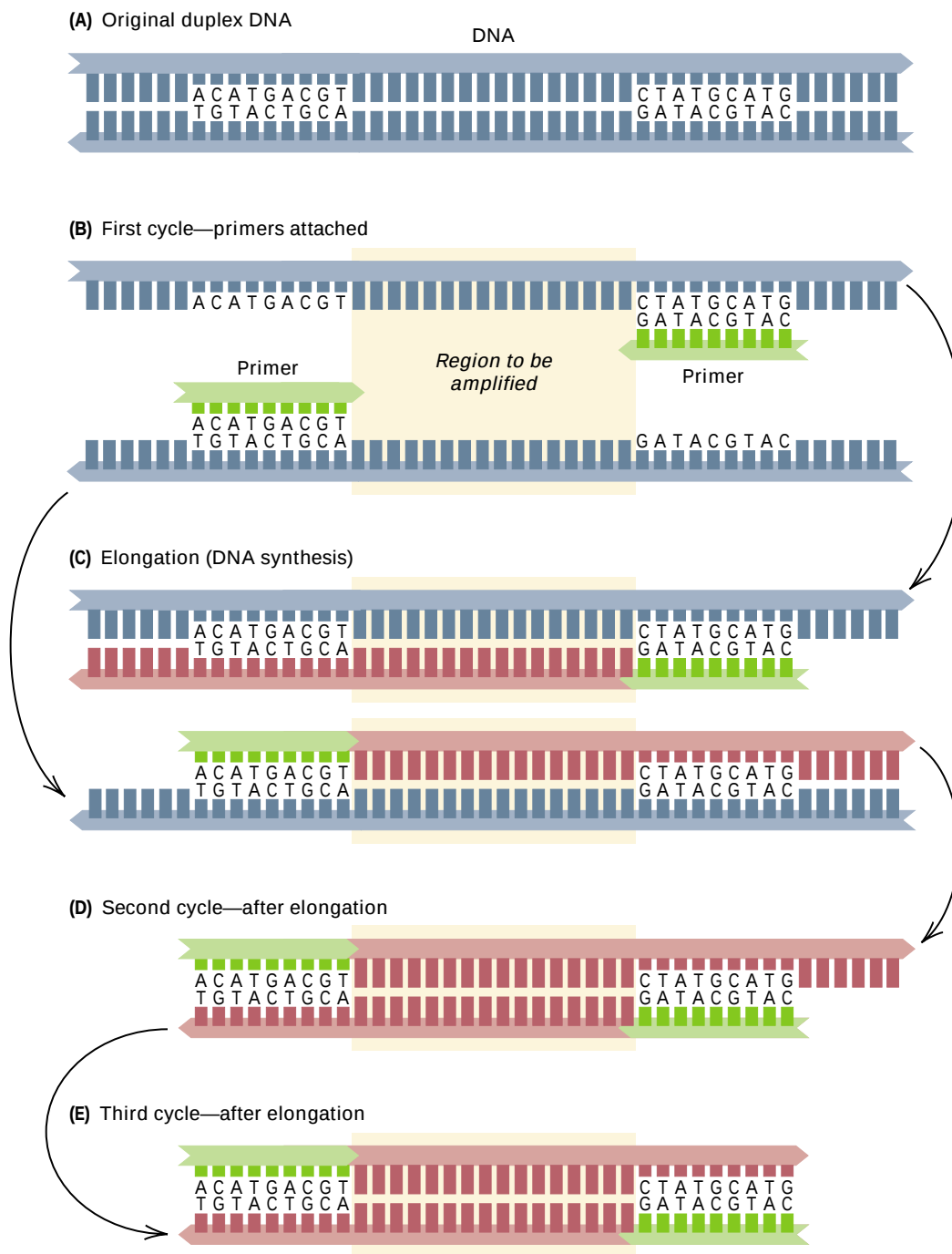


Figure 2.20 Role of primer sequences in PCR amplification. (A) Target DNA duplex (blue), showing sequences chosen as the primer-binding sites flanking the region to be amplified. (B) Primer (green) bound to denatured strands of target DNA. (C) First round of amplification. Newly synthesized DNA is shown in pink. Note that each primer is extended *beyond* the other primer site. (D) Second round of amplification (only one strand shown); in this round, the newly synthesized strand terminates at the opposite primer site. (E) Third round of amplification (only one strand shown); in this round, both strands are truncated at the primer sites. Primer sequences are normally about twice as long as shown here.

When the temperature is lowered again to allow renaturation, the primers, because they are in great excess, become annealed to the separated template strands (part B). Note that the primer sequences are different from each other but complementary to sequences present in opposite strands of the original DNA duplex and flanking the region to be amplified. The primers are oriented with their 3' ends pointing in the direction of the region to be amplified, be-

cause each DNA strand elongates only at the 3' end. After the primers have annealed, each is elongated by DNA polymerase using the original strand as a template, and the newly synthesized DNA strands (red) grow toward each other as synthesis proceeds (part C). Note that:

A region of duplex DNA present in the original reaction mix can be PCR-amplified only if the region is flanked by the primer oligonucleotides.

To start a second cycle of PCR amplification, the temperature is raised again to denature the duplex DNA. Upon lowering of the temperature, the original parental strands anneal with the primers and are replicated as shown in Figure 2.20B and C. The daughter strands produced in the first round of amplification also anneal with primers and are replicated, as shown in part D. In this case, although the daughter duplex molecules are identical in sequence to the original parental molecule, they consist entirely of primer oligonucleotides and nonparental DNA that was synthesized in either the first or the second cycle of PCR. As successive cycles of denaturation, primer annealing, and elongation occur, the original parental strands are diluted out by the proliferation of new daughter strands until eventually, virtually every molecule produced in the PCR has the structure shown in part D.

The power of PCR amplification is that the number of copies of the template strand increases in exponential progression: 1, 2, 4, 8, 16, 32, 64, 128, 256, 512, 1024, and so forth, doubling with each cycle of replication. Starting with a mixture containing as little as one molecule of the fragment of interest, repeated rounds of DNA replication increase the number of amplified molecules exponentially. For example, starting with a single molecule, 25 rounds of DNA replication will result in $2^{25} = 3.4 \times 10^7$ molecules. This number of molecules of the amplified fragment is so much greater than that of the other unamplified molecules in the original mixture that the amplified DNA can often be used without further purification. For example, a single fragment of 3 kb in *E. coli* accounts for only 0.06 percent of the DNA in this organism. However, if this single fragment were replicated through 25 rounds of replication, then 99.995 percent of the resulting mixture would consist of the amplified sequence. A 3-kb fragment of human DNA constitutes only 0.0001 percent of the total genome size. Amplification of a 3-kb fragment of human DNA to 99.995 percent purity would require about 34 cycles of PCR.

An overview of the polymerase chain reaction is shown in Figure 2.21. The DNA sequence to be amplified is again shown in blue and the oligonucleotide primers in green. The oligonucleotides anneal to the

ends of the sequence to be amplified and become the substrates for chain elongation by DNA polymerase. In the first cycle of PCR amplification, the DNA is denatured to separate the strands. The denaturation temperature is usually around 95°C. Then the temperature is decreased to allow annealing in the presence of a vast excess of the primer oligonucleotides. The annealing temperature is typically in the range of 50°C–60°C, depending largely on the G + C content of the oligonucleotide primers. To complete the cycle, the temperature is raised slightly, to about 70°C, for the elongation of each primer. The steps of denaturation, renaturation, and replication are repeated from 20–30 times, and in each cycle the number of molecules of the amplified sequence is doubled.

Implementation of PCR with conventional DNA polymerases is not practical, because at the high temperature necessary for denaturation, the polymerase is itself irreversibly unfolded (denatured) and becomes inactive. However, DNA polymerase isolated from certain organisms is heat stable because the organisms normally live in hot springs at temperatures well above 90°C, such as are found in Yellowstone National Park. Such organisms are said to be **thermophiles**. The most widely used heat-stable DNA polymerase is called *Taq* polymerase, because it was originally isolated from the thermophilic bacterium *Thermus aquaticus*.

PCR amplification is very useful for generating large quantities of a specific DNA sequence. The principal limitation of the technique is that the DNA sequences at the ends of the region to be amplified must be known so that primer oligonucleotides can be synthesized. In addition, sequences longer than about 5000 base pairs cannot be replicated efficiently by conventional PCR procedures, although there are modifications of PCR that allow longer fragments to be amplified. On the other hand, many applications require amplification of relatively small fragments. The major advantage of PCR amplification is that it requires only trace amounts of template DNA. Theoretically only one template molecule is required, but in practice the amplification of a single molecule may fail because the molecule may, by chance, be broken or damaged. But amplification is usually reliable with as

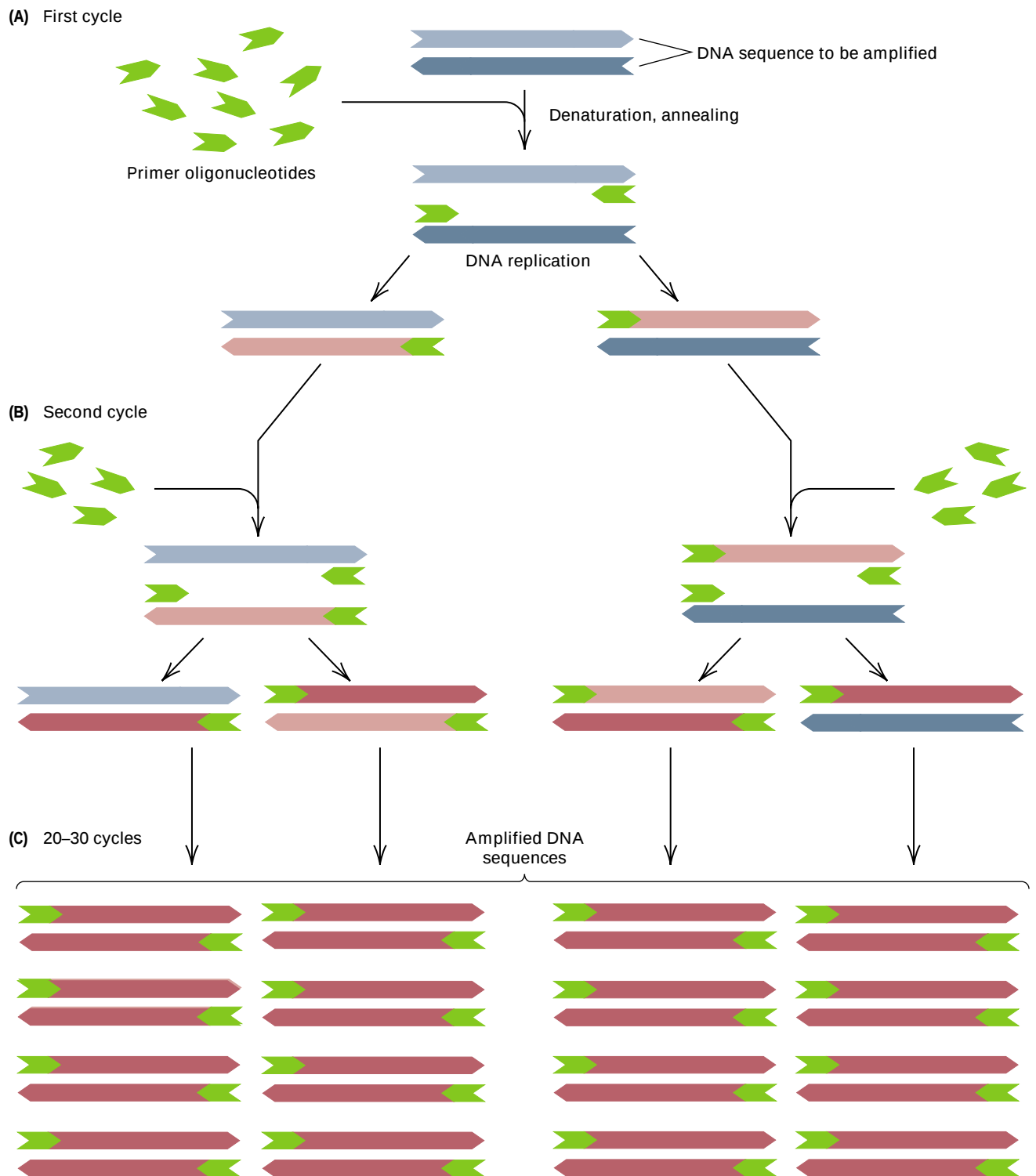


Figure 2.21 Polymerase chain reaction (PCR) for amplification of particular DNA sequences. Only the region to be amplified is shown. Oligonucleotide primers (green) that are complementary to the ends of the target sequence (blue) are used in repeated rounds of denaturation, annealing, and DNA replication. Newly replicated DNA is shown in pink. The number of copies of the target sequence doubles in each round of replication, eventually overwhelming any other sequences that may be present.

few as 10–100 template molecules, which makes PCR amplification 10,000–100,000 times more sensitive than detection via nucleic acid hybridization.

The exquisite sensitivity of PCR amplification has led to its use in DNA typing for criminal cases in which a minuscule amount of biological material has been left behind by the perpetrator (skin cells on a cigarette butt or hair-root cells on a single hair can yield enough template DNA for amplification). In research, PCR is widely used in the study of independent mutations in a gene whose sequence is known in order to identify the molecular basis of each mutation, to study DNA sequence variation among alternative forms of a gene that may be present in natural populations, or to examine differences among genes with the same function in different species. The PCR procedure has also come into widespread use in clinical laboratories for diagnosis. To take just one very important example, the presence of the human immunodeficiency virus (HIV), which causes acquired immune deficiency syndrome (AIDS), can be detected in trace quantities in blood banks via PCR by using primers complementary to sequences in the viral genetic material. These and other applications of PCR are facilitated by the fact that the procedure lends itself to automation by the use of mechanical robots to set up and run the reactions.

2.5 The Terminology of Genetic Analysis

In order to discuss the types of DNA markers that modern geneticists commonly use in genetic analysis, we must first introduce some key terms that provide the essential vocabulary of genetics. These terms can be understood with reference to Figure 2.22. In Chapter 1 we defined a **gene** as an element of heredity, transmitted from parents to offspring in reproduction, that influences one or more hereditary traits. Chemically, a gene is a sequence of nucleotides along a DNA molecule. In a population of organisms, not all copies of a gene may have exactly the same nucleotide sequence. For example, whereas one form of a gene has the codon GCA, which specifies

alanine in the polypeptide chain that the gene encodes, another form of the same gene may have, at the same position, the codon GCG, which also specifies alanine. Hence the two forms of the gene encode the same sequence of amino acids yet differ in DNA sequence. The alternative forms of a gene are called **alleles** of the gene. Different alleles may also code for different amino acid sequences, sometimes with drastic effects. Recall the example of the *PAH* gene for phenylalanine hydroxylase in Chapter 1, in which a change in codon 408 from CGG (arginine) to TGG (tryptophan) results in an inactive enzyme that becomes expressed as the inborn error of metabolism phenylketonuria.

Within a cell, genes are arranged in linear order along microscopic thread-like bodies called **chromosomes**, which we will examine in detail in Chapters 4 and 8. Each human reproductive cell contains one complete set of 23 chromosomes containing 3×10^9 base pairs of DNA. A typical chromosome contains several hundred to several thousand genes. In humans the average is approximately 3500 genes per chromosome. Each chromosome contains a single molecule of duplex DNA along its length, complexed with proteins and very tightly coiled. The DNA in the average human chromosome, when fully extended, has relative dimensions comparable to those of a wet spaghetti noodle 25 miles long; when the DNA is coiled in the form of a chromosome, its physical compaction is comparable to that of the same noodle coiled and packed into an 18-foot canoe.

The physical position of a gene along a chromosome is called the **locus** of the gene. In most higher organisms, including human beings, each cell other than a sperm or egg contains two copies of each type of chromosome—one from the mother and one inherited from the father. Each member of such a pair of chromosomes is said to be **homologous** to the other. (The chromosomes that determine sex are an important exception, discussed in Chapter 4, that we will ignore for now.) At any locus, therefore, each individual carries two alleles, because one allele is present at a corresponding position in each of the homologous maternal and paternal chromosomes (Figure 2.22).

The genetic constitution of an individual is called its **genotype**. For a particular gene, if the two alleles at the locus in an individual are indistinguishable from each other, then for this gene the genotype of the individual is said to be **homozygous** for the allele that is present. If the two alleles at the locus are different from each other, then for this gene the genotype of the individual is said to be **heterozygous** for the alleles that are present. Typographically, genes are indicated in italics, and alleles are typically distinguished by uppercase or lowercase letters (*A* versus *a*), subscripts (A_1 versus A_2), superscripts (a^+ versus a^-), or sometimes just + and -. Using these symbols, homozygous genes would be portrayed by any of these formulas: AA , aa , A_1A_1 , A_2A_2 , a^+a^+ , a^-a^- , $+/+$, or $-/-$. As in the last two examples, the slash is sometimes used to separate alleles present in homologous chromosomes to avoid ambiguity. Heterozygous genes would be portrayed by any of the formulas Aa , A_1A_2 , a^+a^- , or $+/-$. In Figure 2.22, the genotype Bb is heterozygous because the B and b alleles are distinguishable (which is why they are assigned different symbols), whereas the genotype CC is homozygous. These genotypes could also be written as B/b and C/C , respectively.

Whereas the alleles that are present in an individual constitute its genotype, the physical or biochemical expression of the genotype is called the **phenotype**. To put it as simply as possible, the distinction is that the genotype of an individual is what is on the *inside* (the alleles in the DNA), whereas the phenotype is what is on the *outside* (the observable traits, including biochemical traits, behavioral traits, and so forth). The distinction between genotype and phenotype is critically important because there usually is not a one-to-one correspondence between genes and traits. Most complex traits, such as hair color, skin color, height, weight, behavior, life span, and reproductive fitness, are influenced by many genes. Most traits are also influenced more or less strongly by environment. This means that the same genotype can result in different phenotypes, depending on the environment. Compare, for example, two people with a genetic risk for lung cancer; if one smokes and the other does not, the

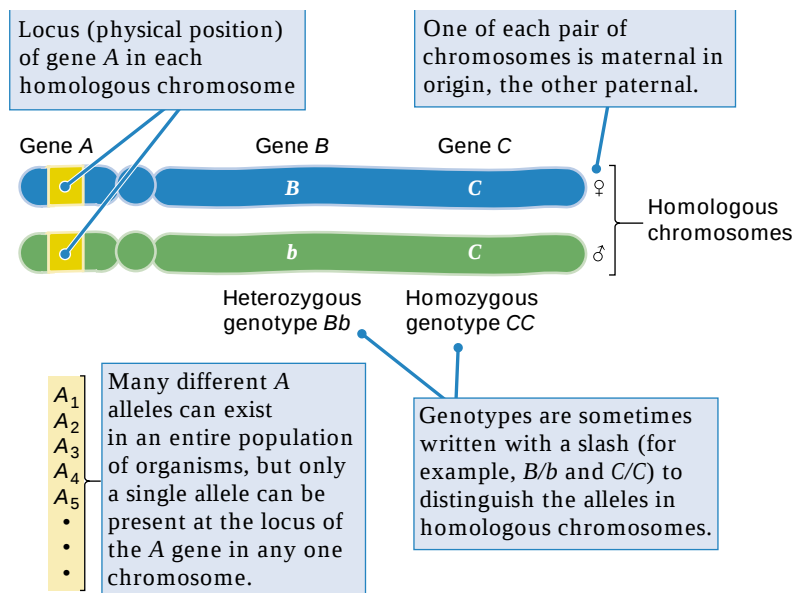


Figure 2.22 Key concepts and terms used in modern genetics. Note that a single gene can have any number of alleles in the population as a whole, but no more than two alleles can be present in any one individual.

smoker is much more likely to develop the disease. Environmental effects also imply that the same phenotype can result from more than one genotype; smoking again provides an example, because most smokers who are not genetically at risk can also develop lung cancer.

2.6 Types of DNA Markers Present in Genomic DNA

Genetic variation, in the form of multiple alleles of many genes, exists in most natural populations of organisms. We have called such genetic differences between individuals *DNA markers*; they are also called **DNA polymorphisms**. (The term *polymorphism* literally means “multiple forms.”) The methods of DNA manipulation examined in Sections 2.3 and 2.4 can be used in a variety of combinations to detect differences among individuals. Anyone who reads the literature in modern genetics will encounter a bewildering variety of acronyms referring to different ways in which genetic polymorphisms are detected. The different approaches are in use because no single

method is ideal for all applications, each method has its own advantages and limitations, and new methods are continually being developed. In this section we examine some of the principal methods for detecting DNA polymorphisms among individuals.

Single-Nucleotide Polymorphisms (SNPs)

A **single-nucleotide polymorphism**, or **SNP** (pronounced “snip”), is present at a particular nucleotide site if the DNA molecules in the population frequently differ in the identity of the nucleotide pair that occupies the site. For example, some DNA molecules may have a T–A base pair at a particular nucleotide site, whereas other DNA molecules in the same population may have a C–G base pair at the same site. This difference constitutes a SNP. The SNP defines two “alleles” for which there could be three genotypes among individuals in the population: homozygous with T–A at the corresponding site in both homologous chromosomes, homozygous with C–G at the corresponding site in both homologous chromosomes, or heterozygous with T–A in one chromosome and C–G in the homologous chromosome. The word *allele* is in quotation marks above because the SNP need not be in a coding sequence, or even in a gene. In the human genome, any two randomly chosen DNA molecules are likely

to differ at one SNP site about every 1000–3000 bp in protein-coding DNA and at about one SNP site every 500–1000 bp in noncoding DNA. Note, in the definition of a SNP, the stipulation that DNA molecules must differ at the nucleotide site “frequently.” This provision excludes rare genetic variation of the sort found in less than 1 percent of the DNA molecules in a population. The reason for the exclusion is that genetic variants that are too rare are not generally as useful in genetic analysis as the more common variants. A catalog of SNPs is regarded as the ultimate compendium of DNA markers, because SNPs are the most common form of genetic differences among people and because they are distributed approximately uniformly along the chromosomes. By the middle of 2001, some 300,000 SNPs are expected to have been identified in human populations and their positions in the chromosomes located.

Restriction Fragment Length Polymorphisms (RFLPs)

Although most SNPs require DNA sequencing to be studied, those that happen to be located within a restriction site can be analyzed with a Southern blot. An example of this situation is shown in Figure 2.23, where the SNP consists of a T–A nucleotide pair in some molecules and a C–G pair in others. In this example, the polymorphic nucleotide

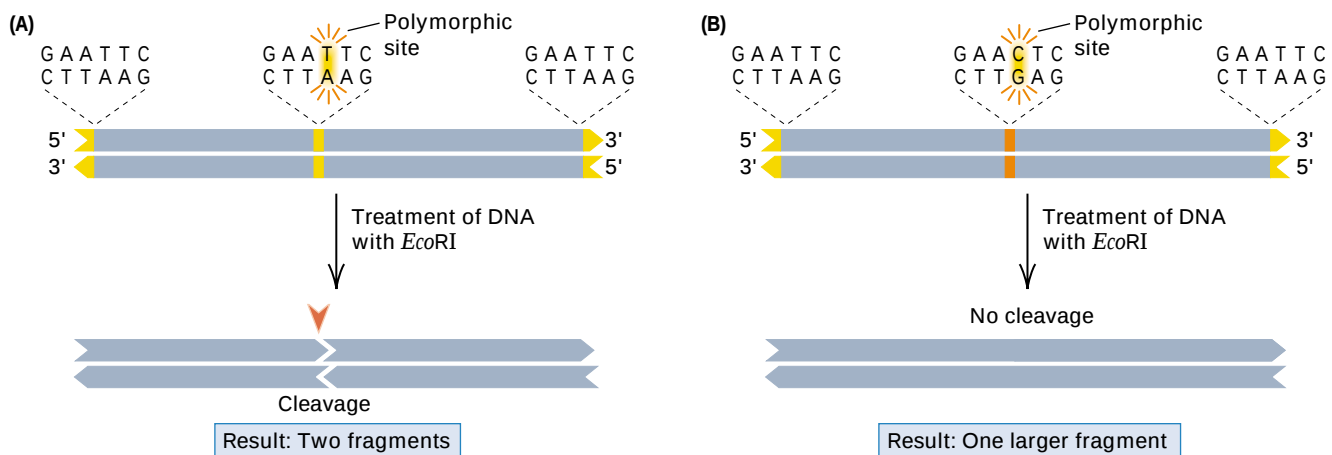
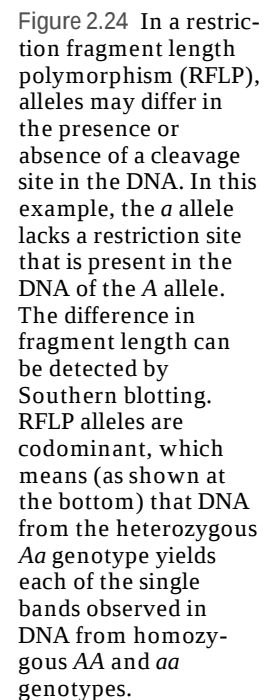


Figure 2.23 A minor difference in the DNA sequence of two molecules can be detected if the difference eliminates a restriction site. (A) This molecule contains three restriction sites for *EcoRI*, including one at each end. It is cleaved into two fragments by the enzyme. (B) This molecule has an altered *EcoRI* site in the middle, in which 5'-GAATTC-3' becomes 5'-GAAGTC-3'. The altered site cannot be cleaved by *EcoRI*, so treatment of this molecule with *EcoRI* results in one larger fragment.

Because RFLPs change the number and size of DNA fragments produced by digestion with a restriction enzyme, they can be detected by the Southern blotting procedure discussed in Section 2.3. An example appears in Figure 2.24. In this case the labeled probe DNA hybridizes near the restriction site at the far left and identifies the position of this restriction fragment in the electrophoresis gel. The duplex molecule labeled “allele A” has a restriction site in the middle and, when cleaved and subjected to electrophoresis, yields a small band that



Origin of the Human Genetic Linkage Map

David Botstein,¹ Raymond L. White,²
Mark Skolnick,³ and Ronald W.
Davis⁴ 1980

¹Massachusetts Institute of Technology,
Cambridge, Massachusetts

²University of Massachusetts Medical
Center, Worcester, Massachusetts

³University of Utah, Salt Lake City, Utah

⁴Stanford University, Stanford,
California

*Construction of a Genetic Linkage Map in
Man Using Restriction Fragment Length
Polymorphisms*

This historic paper stimulated a major international effort to establish a genetic linkage map of the human genome based on DNA polymorphisms. Pedigree studies using these genetic markers soon led to the chromosomal localization and identification of mutant genes for hundreds of human diseases. A more ambitious goal, still only partly achieved, is to understand the genetic and environmental interactions involved in complex traits such as heart disease and cancer. The "small set of large pedigrees" called for in the excerpt was soon established by the Centre d'Etude du Polymorphisme Humain (CEPH) in Paris, France, and made available to investiga-

tors worldwide for genetic linkage studies. Today the CEPH maintains a database on the individuals in these pedigrees that comprises approximately 12,000 polymorphic DNA markers and 2.5 million genotypes.

No method of systematically mapping human genes has been devised, largely because of the paucity of highly polymorphic marker loci. The advent of recombinant DNA technology has suggested a theoretically possible way to define an arbitrarily large number of arbitrarily polymorphic marker loci. . . . A

subset of such polymorphisms can readily be detected as differences in the length of DNA fragments after digestion with DNA sequence-specific restriction endonucleases. These restriction fragment length polymorphisms (RFLPs) can be easily assayed in individuals, facilitating large population studies. . . . [Genetic mapping] of many DNA marker loci should allow the establishment of a set of well-spaced, highly polymorphic genetic markers covering the entire human genome [and enabling] any trait caused wholly or

partially by a major locus segregating in a pedigree to be mapped. Such a procedure would not require any knowledge of the biochemical nature of the trait or of the nature of the alterations in the DNA responsible for the trait. . . .

The most efficient procedure will be to study a small set of large pedigrees which have been genotyped for all known polymorphic markers. . . . The resolution of genetic and environmental components of disease . . . must involve unraveling the underlying

In principle, linked marker loci can allow one to establish, with high certainty, the genotype of an individual.

ing genetic predisposition, understanding the environmental contributions, and understanding the variability of expression of the phenotype. In principle, linked marker loci can allow one to establish, with high certainty, the genotype of an individual and, consequently, assess much more precisely the contribution of modifying factors such as secondary genes, likelihood of expression of the phenotype, and environment.

Source: *American Journal of Human Genetics*
32: 314–331

contains sequences homologous to the probe DNA. The duplex molecule labeled "allele *a*" lacks the middle restriction site and yields a larger band. In this situation there can be three genotypes—AA, Aa, or aa, depending on which alleles are present in the homologous chromosomes—and all three genotypes can be distinguished as shown in the Figure 2.24. Homozygous AA yields only a small fragment, homozygous aa yields only a large fragment, and heterozygous Aa yields both a small and a large fragment. Because the presence of both the A and *a* alleles can be detected in heterozygous Aa genotypes, A and *a* are said to be

codominant. In Figure 2.24, the bands from AA and aa have been shown as somewhat thicker than those from Aa, because each AA genotype has two copies of the A allele and each aa genotype has two copies of the *a* allele, compared with only one copy of each allele in the heterozygous genotype Aa.

Random Amplified Polymorphic DNA (RAPD)

For studying DNA markers, one limitation of Southern blotting is that it requires material (probes available in the form of cloned

DNA) and one limitation of PCR is that it requires sequence information (so primer oligonucleotides can be synthesized). These are not severe handicaps for organisms that are well studied (for example, human beings, domesticated animals and cultivated plants, and model genetic organisms such as yeast, fruit fly, nematode, or mouse), because research materials and sequence information are readily available. But for the vast majority of organisms that biologists study, there are neither research materials nor sequence information. Genetic analysis can still be carried out in these organisms by using an approach called **random amplified polymorphic DNA** or **RAPD** (pronounced “rapid”), described in this section.

RAPD analysis makes use of a set of PCR primers of 8–10 nucleotides whose sequence is essentially *random*. The random primers are tried individually or in pairs in PCR reactions to amplify fragments of genomic DNA from the organism of interest. Because the primers are so short, they

often anneal to genomic DNA at multiple sites. Some primers anneal in the proper orientation and at a suitable distance from each other to support amplification of the unknown sequence between them. Among the set of amplified fragments are ones that can be amplified from some genomic DNA samples but not from others, which means that the presence or absence of the amplified fragment is polymorphic in the population of organisms.

In most organisms it is usually straightforward to identify a large number of RAPDs that can serve as genetic markers for many different kinds of genetic studies. An example of RAPD gel analysis is illustrated in Figure 2.25, where three pairs of primers (sets 1–3) are used to amplify genomic DNA from four individuals in a population. The fragments that amplify are then separated on an electrophoresis gel and visualized after staining with ethidium bromide. Many amplified bands are typically observed for each primer set, but only some of these are

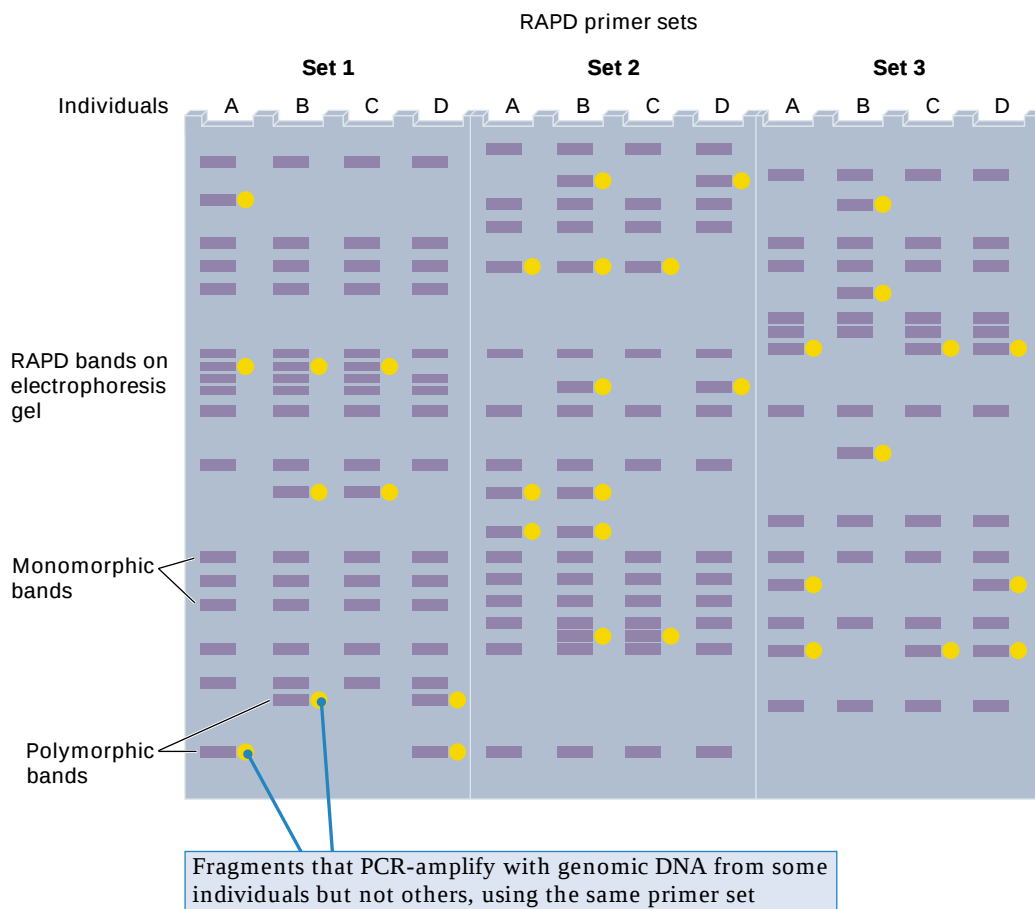


Figure 2.25 Random amplified polymorphic DNA (RAPD) is detected through the use of relatively short primer sequences that, by chance, match genomic DNA at multiple sites that are close enough together to support PCR amplification. Genomic DNA from a single individual typically yields many bands, only some of which are polymorphic in the population. Different sets of primers amplify different fragments of genomic DNA.

polymorphic. These are indicated in Figure 2.25 by the colored dots. The amplified bands that are not polymorphic are said to be **monomorphic** in the sample, which means that they are the same from one individual to the next. This example shows 17 RAPD polymorphisms. Figure 2.26 shows an actual RAPD gel amplified from genomic DNA obtained from small tissue samples from a population of fish (*Camptostoma anomalum*) in the Great Miami River Basin, Ohio. The fish were collected as part of a water quality assessment program to determine whether fish populations in stressful water environments progressively lose their genetic variation (that is, become increasingly monomorphic). Each pair of samples is flanked by a lane containing DNA size standards producing a “ladder” of fragments at 100-bp increments.

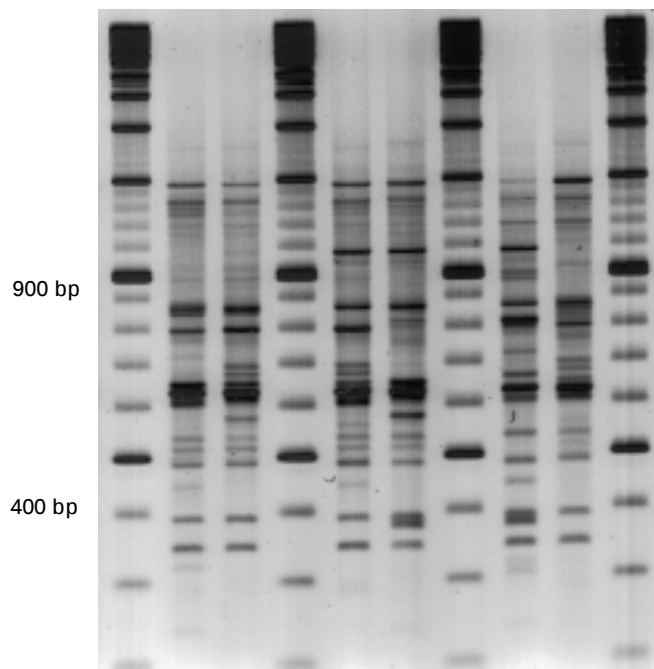


Figure 2.26 RAPD polymorphisms in the stoneroller fish (*Camptostoma anomalum*) trapped in tributaries of the Great Miami River in Ohio. Each pair of samples is flanked by a lane containing DNA size standards; in these lanes, the smallest DNA fragment is 100 base pairs (bp), and each successively larger fragment increases in size by 100 bp. Fragments whose sizes are multiples of 500 bp are present in greater concentration and so yield darker bands. [Courtesy of Michael Simonich, Manju Garg, and Ana Braam (Pathology Associates International, Cincinnati, Ohio).]

Au: Are 400 and 900 base pairs correct? Or is it 100 and 400 base pairs?

Figures 2.24 through 2.26 illustrate an important point:

In modern genetics, the phenotypes that are studied are very often bands in a gel rather than physical or physiological characteristics.

Figure 2.25 offers a good example. Each position at which a band is observed in one or more samples is a phenotype, whether or not the band is polymorphic. For example, primer set 1 yields a total of 19 bands, of which 5 are polymorphic and 14 are monomorphic in the sample. The phenotypes could be named in any convenient way, such as by indicating the primer set and the fragment length. For example, suppose that the smallest amplified fragment for primer set 1 is 125 bp, which is the polymorphic fragment at the bottom left in Figure 2.25. We could name this fragment unambiguously as 1-125 because it is a fragment of 125 bp amplified by primer set 1.

To understand why a DNA band is a *phenotype*, rather than a gene or a genotype, it is useful to assign different names to the “alleles” that do or do not support amplification. (The word *allele* is in quotation marks again, because the 1-125 fragment that is amplified need not be part of a gene.) We are talking only about the 1-125 fragment, so we could call the allele capable of supporting amplification the plus (+) allele and the allele not capable of supporting amplification the minus (−) allele. Then there are three possible genotypes with regard to the amplified fragment: +/+, +/−, and −/−. Using genomic DNA from these genotypes, the homozygous +/+ and heterozygous +/− will both support amplification of the 1-125 fragment, whereas the −/− genotype will not support amplification. Hence the presence of the 1-125 fragment is the *phenotype* observed in both +/+ and +/− genotypes. In other words, with regard to amplification, the + allele is **dominant** to the − allele, because the phenotype (presence of the 1-125 band) is present in both homozygous +/+ and heterozygous +/− genotypes. Therefore, on the basis of the phenotype for the 1-125 band in Figure 2.25, we could say that individuals A and D could have either a +/+ or +/− genotype but that individuals B and C must have genotype −/−.

Amplified Fragment Length Polymorphisms (AFLPs)

Because RAPD primers are small and may not match the template DNA perfectly, the amplified DNA bands often differ a great deal in how dark or light they appear. This variation creates a potential problem, because some exceptionally dark bands may actually result from two amplified DNA fragments of the same size, and some exceptionally light bands may be difficult to visualize consistently. To obtain amplified fragments that yield more uniform band intensities, double-stranded oligonucleotide sequences that match the primer sequences perfectly can be attached to genomic restriction fragments enzymatically prior to amplification. This method, which is outlined in Figure 2.27, yields a class of DNA

polymorphisms known as **amplified fragment length polymorphisms**, or **AFLPs** (usually pronounced by spelling it out). The first step (part A) is to digest genomic DNA with a restriction enzyme; this example uses the enzyme *EcoRI*, whose restriction site is 5'-GAATTC-3'. Digestion yields a large number of restriction fragments flanked by what remains of an *EcoRI* site on each side. In the next step (part B), double-stranded oligonucleotides called **primer adapters**, with single-stranded overhangs complementary to those on the restriction fragments, are ligated onto the restriction fragments using the enzyme *DNA ligase*. The resulting fragments (C) are ready for amplification by means of PCR. Note that the same adapter is ligated onto each end, so a single primer sequence will anneal to both ends and support amplification.

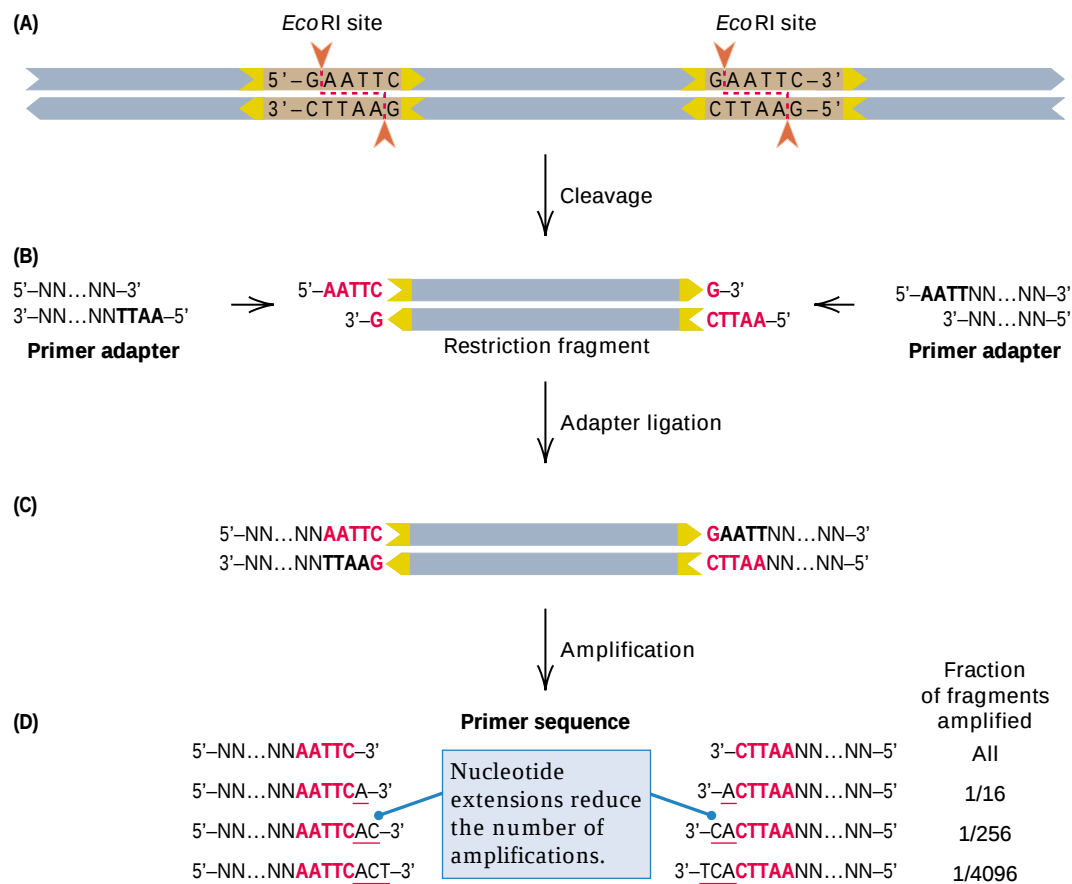


Figure 2.27 An amplified fragment length polymorphism (AFLP). (A and B) Genomic DNA is digested with one or more restriction enzymes (in this case, *EcoRI*). (C) Oligonucleotide adaptors are ligated onto the fragments; note that the single-stranded overhang of the adaptors matches those of the genomic DNA fragments. (D) The resulting fragments are subjected to PCR using primers complementary to the adaptors. The number of amplified fragments can be adjusted by manipulating the number of nucleotides in the adaptors that are also present in the primers.

09131-01-1747P

Photocaptionphotocaptionphotocaptionphotocaptionphotocaptionphotocaptionphotocaptionpho-
tocaptionphotocaptionphotocaptionphotocaption

There are nevertheless a number of choices concerning the primer sequence. A primer that matches the adapters perfectly will amplify all fragments, but this often results in so many amplified fragments that they are not well separated in the gel. A PCR primer must match perfectly at its 3' end to be elongated. Thus, additional nucleotides added to the 3' end reduce the number of amplified fragments, because these primers will amplify only those fragments that, by chance, have a complementary nucleotide immediately adjacent to the *Eco*RI site. For example, the primer sequence in the second row has a single-nucleotide 3' extension; because only $1/4 \times 1/4$ of the fragments would be expected to have a complementary T immediately adjacent to the *Eco*RI site on both sides, this primer is expected to amplify $1/16$ of all the restriction fragments. Similarly, the primer sequence in the third row has a two-nucleotide 3' extension, so this primer is expected to amplify $1/16 \times 1/16 = 1/256$ of all the fragments. One application of AFLP analysis is to organisms with large genomes, such as grasshoppers and crickets, for which RAPD analysis would yield an excessive number of amplified bands. How large is a "large" genome? Compared to the human genome, that of the brown mountain grasshopper *Podisma pedestris* is 7 times larger, and those of North American salamanders in the genus *Amphiuma* are 70 times larger! (Genome size is discussed in Chapter 8.)

Simple Tandem Repeat Polymorphisms (STRPs)

One more type of DNA polymorphism warrants consideration because it is useful in DNA typing for individual identification and for assessing the degree of genetic relatedness between individuals. This type of polymorphism is called a **simple tandem repeat polymorphism (STRP)** because the genetic differences among DNA molecules consist of the number of copies of a short DNA sequence that may be repeated many times in tandem at a particular locus in the genome. STRPs that are present at different loci may differ in the sequence and length of the repeating unit, as well as in the minimum and maximum number of tandem copies that occur in DNA molecules in the population. A STRP with a repeating unit of 2–9 bp is often called a **microsatellite** or a **simple sequence length polymorphism (SSLP)**, whereas a STRP with a repeating unit of 10–60 bp is often called a **minisatellite** or a **variable number of tandem repeats (VNTR)**.

Figure 2.28 shows an example of a STRP with a copy number ranging from 1 through 10. Because the number of copies determines the size of any restriction fragment that includes the STRP, each DNA molecule yields a single-size restriction fragment depending on the number of copies it contains. The STRP in Figure 2.28 has 10 different "alleles" (again, we use quotation marks because the STRP may not

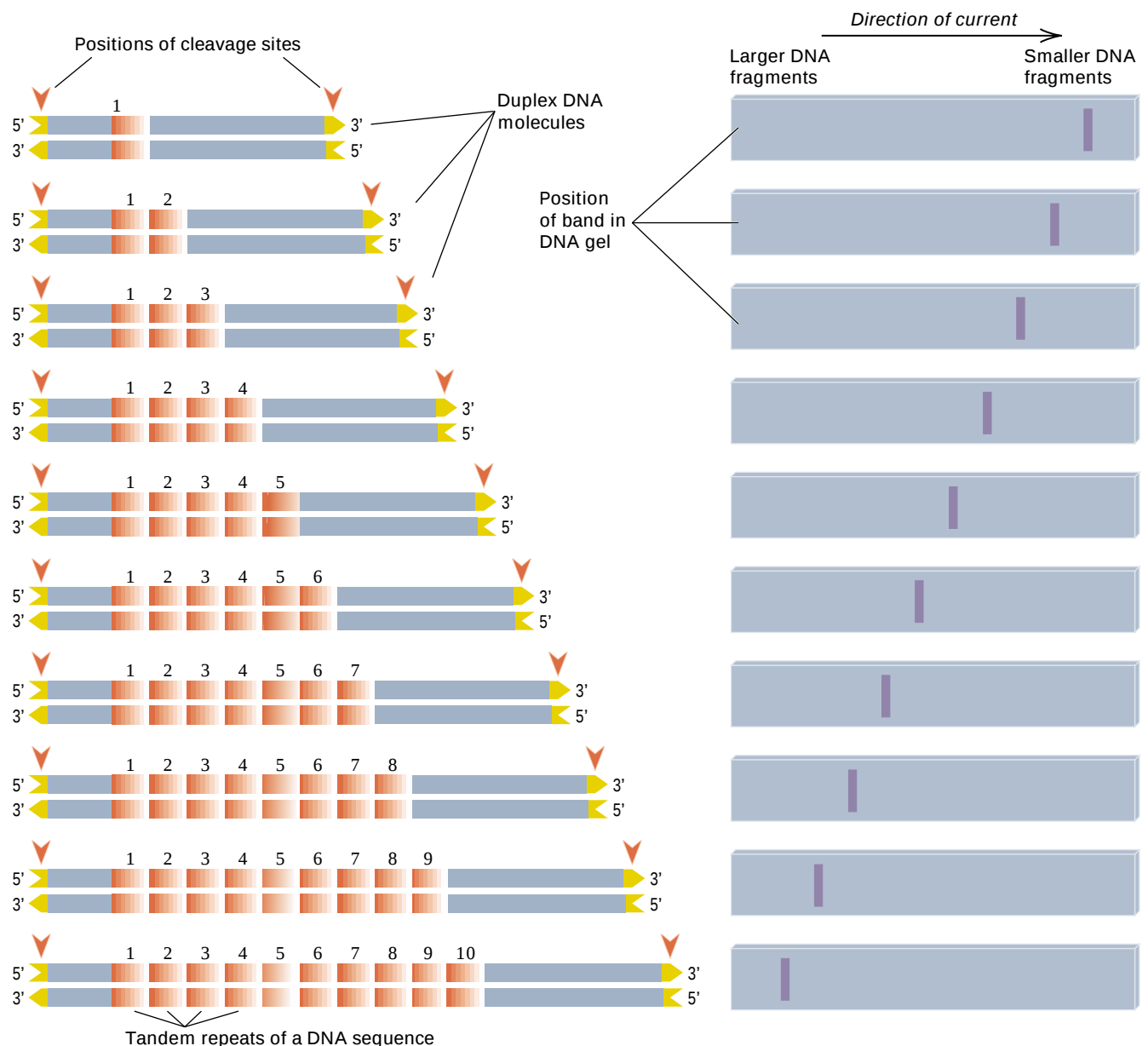


Figure 2.28 In a simple tandem repeat polymorphism (STRP), the alleles in a population differ in the number of copies of a short sequence (typically 2–60 bp) that is repeated in tandem along the DNA molecule. This example shows alleles in which the repeat number varies from 1 to 10. Cleavage at restriction sites flanking the STRP yields a unique fragment length for each allele. The alleles can also be distinguished by the size of the fragment amplified by PCR using primers that flank the STRP.

be in a gene), which could be distinguished either by Southern blotting using a probe to a unique (nonrepeating) sequence within the restriction fragment or by PCR amplification using primers to a unique sequence on either side of the tandem repeats. In this situation the locus is said to have **multiple alleles** in the population. Even with multiple alleles, however, any one chromosome can carry only one of the alleles, and any

individual genotype can carry at most two different alleles. Nevertheless, a large number of alleles means an even larger number of genotypes, which is the feature that gives STRPs their utility in individual identification. For example, even with only 10 alleles in a population of organisms, there could be 10 different homozygous genotypes and 45 different heterozygous genotypes.

More generally, with n alleles there are n homozygous genotypes and $n(n - 1)/2$ heterozygous genotypes, or $n(n + 1)/2$ different genotypes altogether. With STRPs, not only are there a relatively large number of alleles, but no one allele is exceptionally common, so each of the many genotypes in the population has a relatively low frequency. If the genotypes at 6–8 STRP loci are considered simultaneously, then each possible multiple-locus genotype is exceedingly rare. Because of their high degree of variation among people, STRPs are widely used in **DNA typing** (sometimes called **DNA fingerprinting**) to establish individual identity for use in criminal investigations, parentage determinations, and so forth (Chapter 17).

2.7 Applications of DNA Markers

Why are geneticists interested in DNA markers (DNA polymorphisms)? Their interest can be justified on any number of grounds. In this section we consider the reasons most often cited.

Genetic Markers, Genetic Mapping, and “Disease Genes”

Perhaps the key goal in studying DNA polymorphisms in human genetics is to identify the chromosomal location of mutant genes associated with hereditary diseases. In the context of disorders caused by the interaction of multiple genetic and environmental factors, such as heart disease, cancer, diabetes, depression, and so forth, it is important to think of a harmful allele as a **risk factor** for the disease, which increases the probability of occurrence of the disease, rather than as a sole causative agent. This needs to be emphasized, especially because genetic risk factors are often called **disease genes**. For example, the major “disease gene” for breast cancer in women is the gene *BRCA1*. For women who carry a mutant allele of *BRCA1*, the lifetime risk of breast cancer is about 36 percent, and hence most women with this genetic risk factor do *not* develop breast cancer. On the

other hand, among women who are not carriers, the lifetime risk of breast cancer is about 12 percent, and hence many women without the genetic risk factor *do* develop breast cancer. Indeed, *BRCA1* mutations are found in only 16 percent of affected women who have a family history of breast cancer. The importance of a genetic risk factor can be expressed quantitatively as the **relative risk**, which equals the risk of the disease in persons who carry the risk factor as compared to the risk in persons who do not. The relative risk for *BRCA1* equals 3.0 (calculated as 36 percent/12 percent).

The utility of DNA polymorphisms in locating and identifying disease genes results from **genetic linkage**, the tendency for genes that are sufficiently close together in a chromosome to be inherited together. Genetic linkage will be discussed in detail in Chapter 5, but the key concepts are summarized in Figure 2.29, which shows the location of many DNA polymorphisms along a chromosome that also carries a genetic risk factor denoted *D* (for disease gene). Each DNA polymorphism serves as a genetic marker for its own location in the chromosome. The importance of genetic linkage is that DNA markers that are sufficiently close to the disease gene will tend to be inherited together with the disease gene in pedigrees—and the closer the markers, the stronger this association. Hence, the initial approach to the identification of a disease gene is to find DNA markers that are genetically linked with the disease gene in order to identify its chromosomal location, a procedure known as **genetic mapping**. Once the chromosomal position is known, other methods can be used to pinpoint the disease gene itself and to study its functions.

If genetic linkage seems a roundabout way to identify disease genes, consider the alternative. The human genome contains approximately 80,000 genes. If genetic linkage did not exist, then we would have to examine 80,000 DNA polymorphisms, one in each gene, in order to identify a disease gene. But the human genome has only 23 pairs of chromosomes, and because of genetic linkage and the power of genetic mapping, it actually requires only a few hundred DNA polymorphisms to identify the chromosome and approximate location of a genetic risk factor.

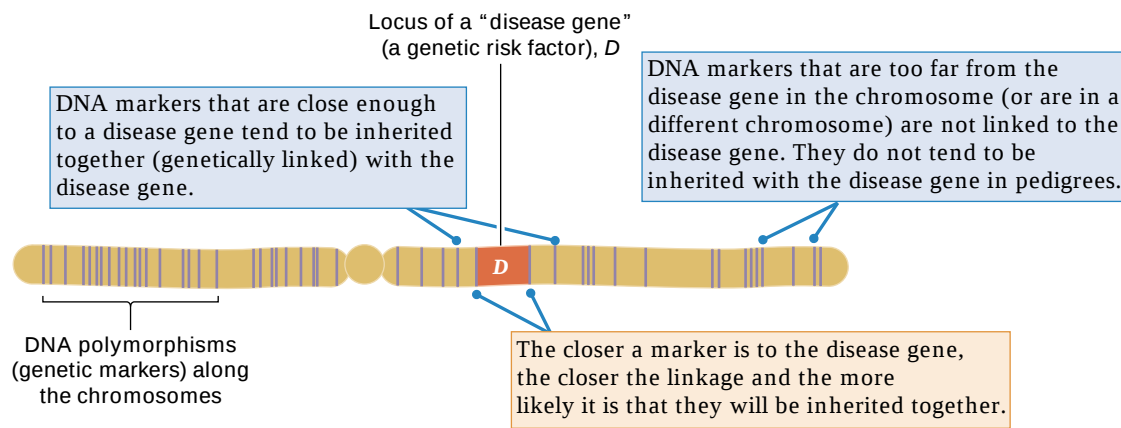


Figure 2.29 Concepts in genetic localization of genetic risk factors for disease. Polymorphic DNA markers (indicated by the vertical lines) that are close to a genetic risk factor (*D*) in the chromosome tend to be inherited together with the disease itself. The genomic location of the risk factor is determined by examining the known genomic locations of the DNA polymorphisms that are linked with it.

Other Uses for DNA Markers

DNA polymorphisms are widely used in all aspects of modern genetics because they provide a large number of easily accessed genetic markers for genetic mapping and other purposes. Among the other uses of DNA polymorphisms are the following.

Individual identification. We have already mentioned that DNA polymorphisms have application as a means of DNA typing (DNA fingerprinting) to identify different individuals in a population. DNA typing in other organisms is used to determine individual animals in endangered species and to identify the degree of genetic relatedness among individual organisms that live in packs or herds. For example, DNA typing in wild horses has shown that the wild stallion in charge of a harem of mares actually sires fewer than one-third of the foals.

Epidemiology and food safety science. DNA typing also has important applications in tracking the spread of viral and bacterial epidemic diseases, as well as in identifying the source of contamination in contaminated foods.

Human population history. DNA polymorphisms are widely used in anthropology to reconstruct the evolutionary origin, global expansion, and diversification of the human population.

Improvement of domesticated plants and animals. Plant and animal breeders have turned to DNA polymorphisms as genetic markers in pedigree studies to identify, by genetic mapping, genes that are associated with favorable traits in order to incorporate these genes into currently used varieties of plants and breeds of animals.

History of domestication. Plant and animal breeders also study genetic polymorphisms to identify the wild ancestors of cultivated plants and domesticated animals, as well as to infer the practices of artificial selection that led to genetic changes in these species during domestication.

DNA polymorphisms as ecological indicators. DNA polymorphisms are being evaluated as biological indicators of genetic diversity in key indicator species present in biological communities exposed to chemical, biological, or physical stress. They are also used to monitor genetic diversity in endangered species and species bred in captivity.

Evolutionary genetics. DNA polymorphisms are studied in an effort to describe the patterns in which different types of genetic variation occur throughout the genome, to infer the evolutionary mechanisms by which genetic variation is maintained, and to illuminate the processes by which genetic polymorphisms within

species become transformed into genetic differences between species.

Population studies. Population ecologists employ DNA polymorphisms to assess the level of genetic variation in diverse populations of organisms that differ in genetic organization (prokaryotes, eukaryotes, organelles), population size, breeding structure, or life-history characters, and they use genetic polymorphisms within subpopulations of a species as indicators of

population history, patterns of migration, and so forth.

Evolutionary relationships among species. Differences in homologous DNA sequences between species is the basis of molecular systematics, in which the sequences are analyzed to determine the ancestral history (phylogeny) of the species and to trace the origin of morphological, behavioral, and other types of adaptations that have arisen in the course of evolution.

Chapter Summary

The sequence of bases in the human genome is 99.9% identical from one person to the next. The remaining 0.1%—comprising 3 million base pairs—differs among individuals. Included in these differences are many mutations that cause or increase the risk of disease, but the majority of the differences are harmless in themselves. Any of these differences between genomes can be used as a genetic marker. Genetic markers are widely employed in genetics to serve as positional landmarks along a chromosome or to identify particular cloned DNA fragments. The manipulation of DNA molecules to identify genetic markers is the basic experimental operation in modern genetics.

A DNA strand is a polymer of deoxyribonucleotides, each composed of a nitrogenous base, a deoxyribose sugar, and a phosphate. Sugars and phosphates alternate in forming a polynucleotide chain with one terminal 3'-OH group and one terminal 5'-P group. In double-stranded (duplex) DNA, the two strands are paired and antiparallel. Each end of the double helix carries a terminal 3'-OH group in one strand and a terminal 5'-P group in the other strand. The four bases found in DNA are the purines, adenine (A) and guanine (G), and the pyrimidines, cytosine (C) and thymine (T). Equal numbers of purines and pyrimidines are found in double-stranded DNA (Chargaff's rules), because the bases are paired as A-T pairs and G-C pairs. The hydrogen-bonded base pairs, along with hydrophobic base stacking of the nucleotide pairs in the core of the double helix, hold the two polynucleotide strands together in a double helix.

Duplex DNA can be cleaved into fragments of defined length by restriction enzymes, each of which cleaves DNA at a specific recognition sequence (restriction site) usually four or six nucleotide pairs in length. These fragments can be separated by electrophoresis. The positions of particular restriction fragments in a gel can be visualized by means of nucleic acid hybridization, in which strands of duplex DNA that have been separated (denatured) by heating are mixed and come together (renature) with strands having complementary nucleotide sequences. In a Southern blot, denatured and labeled probe DNA is mixed with denatured DNA made up of restriction fragments that have been transferred to a filter membrane after electrophoresis. The probe DNA anneals and forms

stable duplexes with whatever fragments contain sufficiently complementary base sequences, and the positions of these duplexes can be determined by exposing the filter to x-ray film on which radioactive emission (or, in some procedures, light emission) produces an image of the band. Particular DNA sequences can also be amplified without cloning by means of the polymerase chain reaction (PCR), in which short, synthetic oligonucleotides are used as primers to replicate repeatedly and amplify the sequence between them. The primers must flank, and have their 3' ends oriented toward, the region to be amplified, because DNA polymerase can elongate the primers only by the addition of successive nucleotides to the 3' end of the growing chain. Each round of PCR amplification results in a doubling of the number of amplified fragments.

Most genes are present in pairs in the nonreproductive cells of most animals and higher plants. One member of each gene pair is in the chromosome inherited from the maternal parent, and the other member of the gene pair is at a corresponding location (locus) in the homologous chromosome inherited from the paternal parent. A gene can have different forms that correspond to differences in DNA sequence. The different forms of a gene are called alleles. The particular combination of alleles present in an organism constitutes its genotype. The observable characteristics of the organism constitute its phenotype. In an organism, if the two alleles of a gene pair are the same (for example, AA or aa), then the genotype is homozygous for the A or a allele; if the alleles are different (Aa), then the genotype is heterozygous. Even though each genotype can include at most two alleles, multiple alleles are often encountered among the individuals in natural populations.

DNA polymorphisms (DNA markers) are common in natural populations of most organisms. Among the most widely used DNA polymorphisms are single-nucleotide polymorphisms (SNPs), restriction fragment length polymorphisms (detected by Southern blots), and such PCR-based polymorphisms as random amplified polymorphic DNA (RAPD), amplified fragment length polymorphisms (AFLPs), and simple tandem repeat polymorphisms (STRPs). DNA polymorphisms are used in genetic map-

ping studies to identify DNA markers that are genetically linked to disease genes (genetic risk factors) in the chromosome in order to pinpoint their location. They are also used in DNA typing for identifying individuals, tracking the course of virus and bacterial epidemics, studying hu-

man population history, and improving cultivated plants and domesticated animals, as well as for the genetic monitoring of endangered species and for many other purposes.

Key Terms

allele	genomic DNA	probe
amplification	genotype	purine
amplified fragment length polymorphism (AFLP)	heterozygous	pyrimidine
antiparallel	homologous chromosomes	random amplified polymorphic DNA (RAPD)
band	homozygous	relative risk
base composition	hydrogen bond	renaturation
base pairing	hydrophobic interaction	restriction endonuclease
base stacking	kilobase (kb)	restriction enzyme
B form of DNA	locus	restriction fragment
blunt end	major groove	restriction fragment length polymorphism (RFLP)
Chargaff's rules	microsatellite	restriction map
chromosome	minisatellite	restriction site
codominant	minor groove	risk factor
denaturation	monomorphic	simple sequence length polymorphism (SSLP)
denatured DNA	multiple alleles	simple tandem repeat polymorphism (STRP)
disease gene	nucleic acid hybridization	single-nucleotide polymorphism (SNP)
DNA cloning	nucleoside	Southern blot
DNA fingerprinting	nucleotide	sticky end
DNA marker	oligonucleotide	thermophile
DNA polymerase	palindrome	3'-OH (hydroxyl) group
DNA polymorphism	percent G + C	variable number of tandem repeats (VNTR).
DNA typing	phenotype	Z form of DNA
dominant	phosphodiester bond	
5'-P (phosphate) group	polarity	
gel electrophoresis	polymerase chain reaction (PCR)	
gene	polynucleotide chain	
genetic linkage	primer	
genetic mapping	primer adapter	
genetic marker		

Review the Basics

- What four bases are commonly found in the nucleotides in DNA? Which form base pairs?
- Which chemical groups are present at the extreme 3' and 5' ends of a single polynucleotide strand?
- What does it mean to say that a single strand of DNA has a polarity? What does it mean to say that the DNA strands in a duplex molecule are antiparallel?
- What are restriction enzymes and why are they important in the study of particular DNA fragments? What does it mean to say that most restriction sites are palindromes?
- Describe how a Southern blot is carried out. Explain what it is used for. What is the role of the probe?
- How does the polymerase chain reaction work? What is it used for? What information about the target sequence must be known in advance? What is the role of the oligonucleotide primers?
- What is a DNA marker? Explain how harmless DNA markers can serve as aids in identifying disease genes through genetic mapping.
- Define and give an example of each of the following key genetic terms: locus, allele, genotype, heterozygous, homozygous, phenotype.

GeNETics on the Web will introduce you to some of the most important sites for finding genetic information on the Internet. To explore these sites, visit the Jones and Bartlett home page at

<http://www.jbpub.com/genetics>

For the book *Genetics: Analysis of Genes and Genomes*, choose the link that says Enter GeNETics on the Web. You will be presented with a chapter-by-chapter list of highlighted keywords. Select any highlighted keyword and you will be linked to a Web site containing genetic information related to the keyword.

- DNA is like Coca-Cola? According to this keyword site, it is. It contains sugar, which is highly soluble in water; phosphate groups, which are of moderate solubility; and bases, which have extremely low solubility. (The base in the soft drink is caffeine, which is chemically similar to adenine and can sometimes be incorporated into DNA, causing a mutation.) As this keyword site emphasizes, the most important property

contributing to the stability of double-stranded DNA is the stacking of the base pairs on top of one another as a result of hydrophobic interactions. For further discussion of this feature of DNA, and much else of interest regarding the discovery and analysis of this critical biological macromolecule, consult the keyword site.

- The concept of the polymerase chain reaction (PCR) occurred to Kary Mullis one night while cruising on Route 128 from San Francisco to Mendocino. He immediately realized that this approach would be unique in its ability to amplify, at an exponential rate, a specific nucleotide sequence present in a vanishingly small quantity amid a much larger background of total nucleic acid. Once its feasibility was demonstrated, PCR was quickly recognized as a major technical advance in molecular biology. The new technique earned Mullis the 1993 Nobel Prize in chemistry, and today it is the basis of a large number of experimental and diagnostic procedures. At this keyword site you can learn more about the

Guide to Problem Solving

Problem 1 Distinguish between base pairing and base stacking in double-stranded DNA.

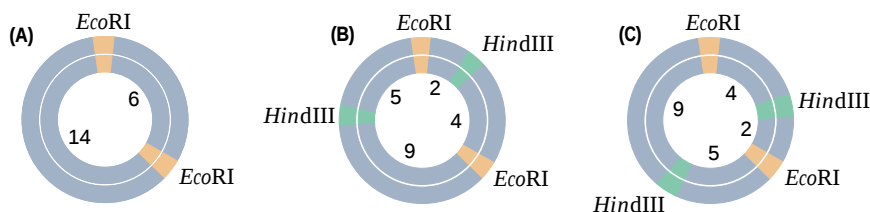
Answer Base pairing is the hydrogen bonding between corresponding bases in opposite strands of duplex DNA; A (adenine) is paired with T (thymine), and G (guanine) is paired with C (cytosine). Base stacking refers to the hydrophobic (water-hating) interaction between consecutive base pairs along a DNA duplex, which promotes the formation of a “stack” of base pairs with the sugar-phosphate backbones of the strands running along outside.

Problem 2 The restriction enzyme *EcoRI* cleaves double-stranded DNA at the sequence 5'-GAATTC-3', and the restriction enzyme *HindIII* cleaves at 5'-AAGCTT-3'. A 20-kilobase (kb) circular plasmid is digested with each enzyme individually and then in combination, and the resulting fragment sizes are determined by means of electrophoresis. The results are as follows:

<i>EcoRI</i> alone	fragments of 6 kb and 14 kb
<i>HindIII</i> alone	fragments of 7 kb and 13 kb
<i>EcoRI</i> and <i>HindIII</i>	fragments of 2 kb, 4 kb, 5 kb and 9 kb

How many possible restriction maps are compatible with these data? For each possible restriction map, make a diagram of the circular molecule and indicate the relative positions of the *EcoRI* and *HindIII* restriction sites.

Answer Because the single-enzyme digests give two bands each, there must be two restriction sites for each enzyme in the molecule. Furthermore, because digestion with *HindIII* makes both the 6-kb and the 14-kb restriction fragments disappear, each of these fragments must contain one *HindIII* site. Considering the sizes of the fragments in the double digest, the 6-kb *EcoRI* fragment must be cleaved into 2-kb and 4-kb fragments, and the 14-kb *EcoRI* fragment must be cleaved into 5-kb and 9-kb fragments. Two restriction maps are compatible with the data, depending on which end of the 6-kb *EcoRI* fragment the *HindIII* site is nearest. The position of the remaining *HindIII* site is determined by the fact that the 2-kb and 5-kb fragments in the double digest must be adjacent in the intact molecule in order for a 13-kb fragment to be produced by *HindIII* digestion alone. The accompanying figure shows the relative positions of the



development of PCR from Mullis's original conception, including two major innovations that were necessary to perfect the process.

- Human beings rely on plants for food, shelter, and medicines. Although at least 5000 species are cultivated, modern agricultural research emphasizes a few widely cultivated crops while largely ignoring plants such as Bambara groundnut (*Vigna subterranea*), breadfruit (*Artocarpus altilis*), carob (*Ceratonia siliqua*), coriander (*Coriandrum sativum*), emmer wheat (*Triticum dicoccum*), oca (*Oxalis tuberosa*), and ulluco (*Ullucus tuberosus*). To learn more about these minor crops and the use of molecular markers for characterizing and preserving their genetic diversity, consult this keyword site.

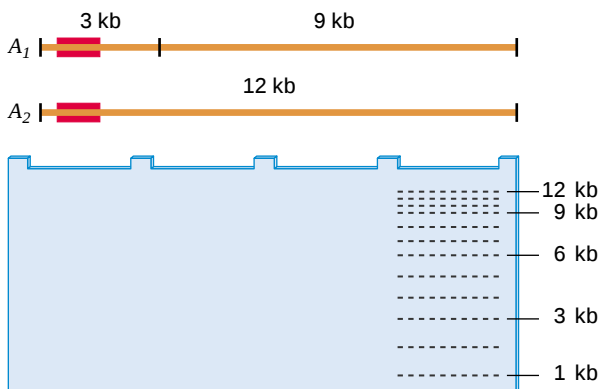
- The Mutable Site changes frequently. Each new update includes a different site that highlights genetics resources available on the World Wide Web. Select the Mutable Site for Chapter 2 and you will be linked automatically.

- The Pic Site showcases some of the most visually appealing genetics sites on the World Wide Web. To visit the genetics Web site pictured below, select the PIC Site for Chapter 2.



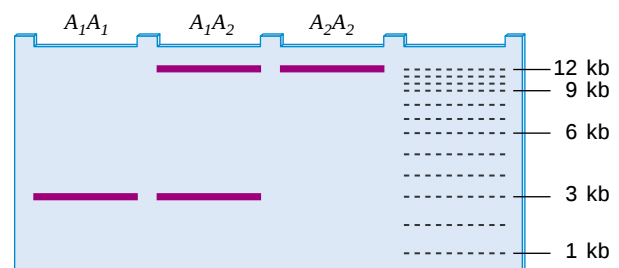
EcoRI sites (part A). Parts B and C are the two possible restriction maps, which differ according to whether the *EcoRI* site at the top generates the 2-kb or the 4-kb fragment in the double digest.

Problem 3 The accompanying diagram shows the positions of restriction sites (tick marks) for a particular restriction enzyme that can be present in the DNA at a locus in a human chromosome. The DNA present in any particular chromosome may be that shown at the top or that shown at the bottom. A probe DNA binds to the fragments at the position shown by the rectangle. With respect to an RFLP based on these fragments, three genotypes are possible. What are they? Use the symbol A_1 to refer to the allele that yields the upper DNA fragment, and use A_2 to refer to the allele that yields the lower DNA fragment. In the



accompanying gel diagram, indicate the genotypes across the top and the phenotype (band position or positions) expected for each genotype. The scale on the right shows the expected positions of fragments ranging in size from 1 to 12 kb.

Answer After cleavage with the restriction enzyme, the A_1 -type DNA yields a 3-kb fragment that binds with the probe (yielding a 3-kb band) and a 9-kb fragment that does not bind with the probe (yielding no visible band), whereas the A_2 -type DNA yields a 12-kb fragment that binds with the probe (yielding a 12-kb band). A particular chromosome may carry allele A_1 or allele A_2 . Because individuals have two copies of each chromosome (except for the sex chromosomes), any individual may carry A_1A_1 , A_1A_2 , or A_2A_2 . DNA from homozygous A_1A_1 genotypes yields a 3-kb band, that from heterozygous A_1A_2 genotypes yields both a 3-kb and a 12-kb band, and that from homozygous A_2A_2 genotypes yields a 12-kb band. The expected phenotypes are illustrated here.



5'-GTACGGGCAATGGTAATTTTTCAGGAACCAGGGCCCTTAAGCCGTC-3'
 3'-CATGCCCGTTACCATTAATAAACTCCTTGGTCCCAGGAATTCGGCAG-5'

Problem 4

Problem 4 A geneticist plans to use the polymerase chain reaction (PCR) to amplify part of the DNA sequence shown below, using oligonucleotide primers that are hexamers matching the regions shown in red. (In practice, hexamers are too short for most purposes.) State the sequence of the primer oligonucleotides that should be used, including the polarity, and give the sequence of the DNA molecule that results from amplification.

Answer The primers must be able to base-pair with the chosen primer sites and must be oriented with their 3' ends facing one another. Thus the “forward primer” (the

one that is elongated in a left-to-right direction) should have the sequence 5'-GCAATG-3' and the “reverse primer” (the one that is elongated in a right-to-left direction) should have the sequence 3'-TTCGGC-5'. The resulting amplified sequence is shown below.

5'-GCAATGGTAATTTTTCAGGAACCAGGGCCCTTAAGCCG-3'
 3'-CGTTACCATTAATAAACTCCTTGGTCCCAGGAATTCGGC-5'

AU/ED: Space in art for prob 2.6 reduced slightly to make page.OK?

Analysis and Applications

2.1 Many restriction enzymes produce restriction fragments that have “sticky ends.” What does this mean?

2.2 Which of the following sequences are palindromes and which are not? Explain your answer. Symbols such as (A/T) mean that the site may be occupied by (in this case) either A or T, and N stands for any nucleotide.

- 5'-AATT-3'
- 5'-AAAA-3'
- 5'-AANTT-3'
- 5'-AA(A/T)AA-3'
- 5'-AA(G/C)TT-3'

2.3 The following list gives half of each of a set of palindromic restriction sites. What is the complete sequence of each restriction site? (N stands for any nucleotide.)

- 5'-AA??-3'
- 5'-ATG??-3'
- 5'-GGN??-3'
- 5'-ATNN??-3'

2.4 Apart from the base sequence, what is different about the ends of restriction fragments produced by the following restriction enzymes? (The downward arrow represents the site of cleavage in each strand.)

- Hae*III (5'-GG ↓ CC-3')
- Mae*I (5'-C ↓ TAG-3')
- Cfo*I (5'-GCG ↓ C-3')

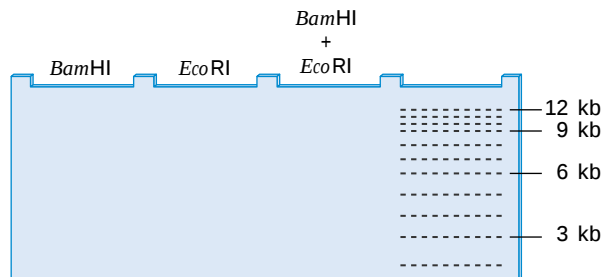
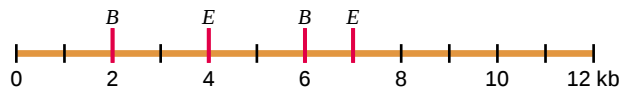
2.5 A solution contains double-stranded DNA fragments of size 3 kb, 6 kb, 9 kb, and 12 kb. They are separated in an electrophoresis gel. In the diagram of the gel at the right, match the fragment sizes with the correct bands.



2.6 The linear DNA fragment shown here has cleavage sites for *Bam*HI (B) and *Eco*RI (E). In the accompanying diagram of an electrophoresis gel, indicate the positions at which bands would be found after digestion with:

- Bam*HI alone
- Eco*RI alone
- Bam*HI and *Eco*RI together

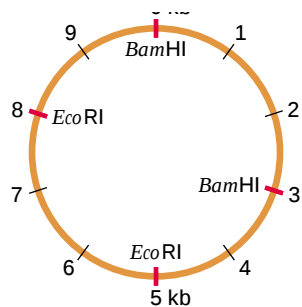
The dashed lines on the right indicate the positions to which bands of 1–12 kb would migrate.



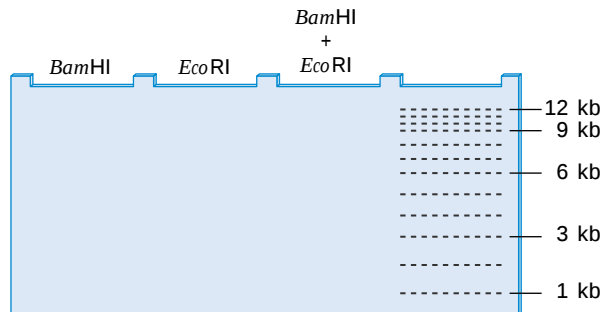
2.7 The circular DNA molecule shown at the top of page 81 has cleavage sites for *Bam*HI and *Eco*RI. In the accompanying diagram of an electrophoresis gel, indicate the positions at which bands would be found after digestion with:

- Bam*HI alone
- Eco*RI alone
- Bam*HI and *Eco*RI together

The dashed lines on the right indicate the positions to which bands of 1–12 kb would migrate.

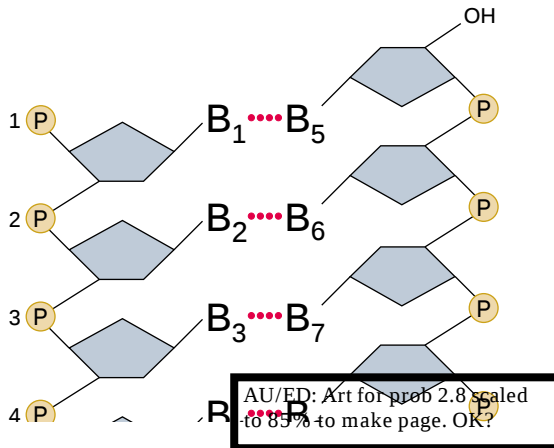


AU/ED: Space in art for prob 2.7 reduced slightly to make page.OK?



2.8 Consider the accompanying diagram of a region of duplex DNA, in which the B's represent bases in Watson-Crick pairs. Specify as precisely as possible the identity of:

- B_5 , assuming that $B_1 = A$
- B_6 , assuming that $B_2 = C$
- B_7 , assuming that $B_3 = \text{purine}$
- B_8 , assuming that $B_4 = A$ or T



AU/ED: Art for prob 2.8 scaled to 85% to make page. OK?

2.9 Refer to the DNA molecule diagrammed in Problem 2.8. In the precursor nucleotides of this molecule, which base was each of the phosphate groups 1–4 associated with?

2.10 In a random sequence consisting of equal proportions of all four nucleotides, what is the probability that a particular short sequence of nucleotides matches a restriction site for:

- A restriction enzyme with a 4-base cleavage site?
- A restriction enzyme with a 6-base cleavage site?
- A restriction enzyme with an 8-base cleavage site?

2.11 In a random sequence consisting of equal proportions of all four nucleotides, what is the average distance between restriction sites for:

- A restriction enzyme with a 4-base cleavage site?
- A restriction enzyme with a 6-base cleavage site?
- A restriction enzyme with an 8-base cleavage site?

2.12 If human DNA were essentially a random sequence of 3×10^9 bp with equal proportions of all four nucleotides (this is an oversimplification), approximately how many restriction fragments would be expected from cleavage with

- A “4-cutter” restriction enzyme?
- A “6-cutter” restriction enzyme?
- An “8-cutter” restriction enzyme?

2.13 Consider the restriction enzymes *Bam*HI (cleavage site 5'-G ↓ GATCC-3') and *Sau*3A (cleavage site 5'-↓ GATC-3'), where the downward arrow denotes the site of cleavage in each strand. Is every *Bam*HI site a *Sau*3A site? Is every *Sau*3A site a *Bam*HI site? Explain your answer.

2.14 A DNA duplex with the sequence shown here is cleaved with *Bam*HI (cleavage site 5'-G ↓ GATCC-3'), where the arrow denotes the site of cleavage in each strand. If the resulting fragments were brought together in the right order, and the breaks in the backbones repaired, what possible DNA duplexes would be expected?

5'-ATTGGATCCAAACCCCAAAGGATCCTTA-3'
3'-TAACCTAGGTTTGGGGTTTCTAGGAAT-5'

2.15 The restriction enzymes *Pst*I, *Pvu*II, and *Mlu*I have the following restriction sites, where the arrow indicates the site of cleavage in each strand.

*Pst*I 5'-CTGCA ↓ G-3'
*Pvu*II 5'-CAG ↓ CTG-3'
*Mlu*I 5'-A ↓ CGCGT-3'

A DNA duplex with the sequence above is digested. What fragments would result from cleavage with:

- Pst*I?
- Pvu*II?
- Mlu*I?

2.16 With regard to the restriction enzymes and the DNA duplex in Problem 2.15, what fragments would result from digestion with

- Pst*I and *Mlu*I?
- Pvu*II and *Mlu*I?

Problems 2.15, 2.16

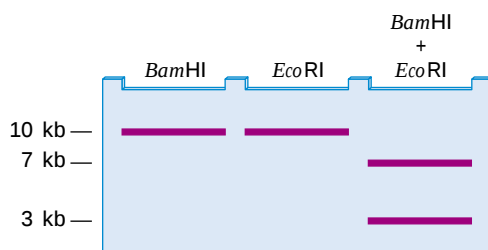
5'-ATGCCCTGCGATACCATGACGCGTTACGAGCTGATCGAAACGCGTATATATGCC-3'
3'-TACGGGACGTCATGGTACTGCGCAATGCGTCGACTAGCTTTGCGCATATATACGG-5'

2.17 Consider the sequence:

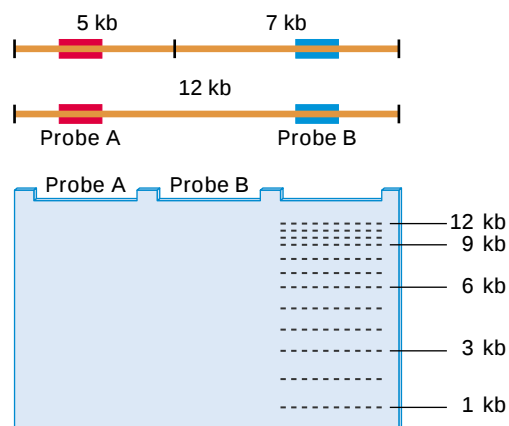
5'-CTGCAGGTG-3'
3'-GACGTCCAC-5'

If this sequence were cleaved with *Pst*I (5'-CTGCA ↓ G-3'), could it still be cleaved with *Pvu*II (5'-CAG ↓ CTG-3')? If it were cleaved with *Pvu*II, could it still be cleaved with *Pst*I? Explain your answer.

2.18 A circular DNA molecule is cleaved with *Bam*HI, *Eco*RI, or the two restriction enzymes together. The accompanying diagram shows the resulting electrophoresis gel, with the band sizes indicated. Draw a diagram of the circular DNA, showing the relative positions of the *Bam*HI and *Eco*RI sites.

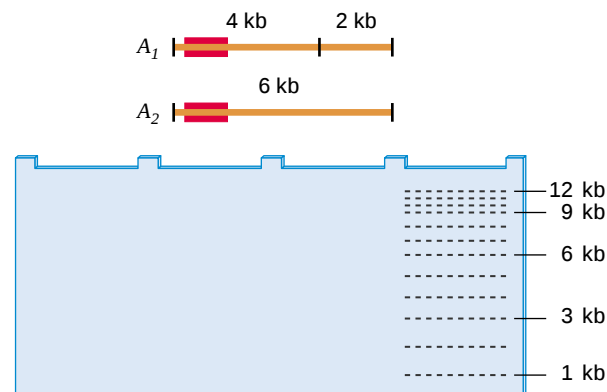


2.19 In the diagrams of DNA fragments shown here, the tick marks indicate the positions of restriction sites for a particular restriction enzyme. A mixture of the two types of molecules is digested and analyzed with a Southern blot using either probe A or probe B, which hybridizes to the fragments where shown by the rectangles. In the accompanying gel diagram, indicate the bands that would result from the use of each of these probes. (The scale on the right shows the expected positions of fragments from 1 to 12 kb.)

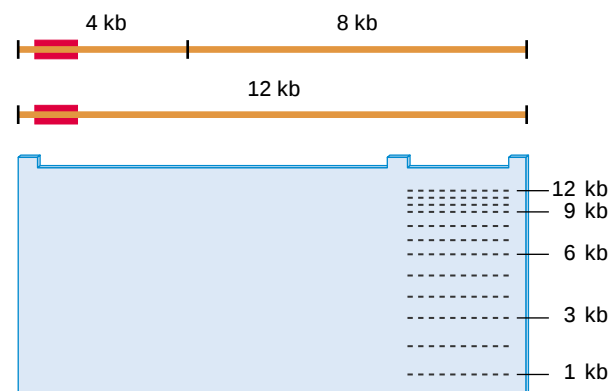


2.20 In the accompanying diagram, the tick marks indicate the positions of restriction sites in two alternative DNA fragments that can be present at the A locus in a human chromosome. An RFLP analysis is carried out, using probe DNA that binds to the fragments at the position shown by the rectangle. With respect to this RFLP, how many genotypes are possible? (Use the symbol A_1 to refer to the allele that yields the upper DNA fragment, and use A_2 to refer to the allele that yields the lower DNA frag-

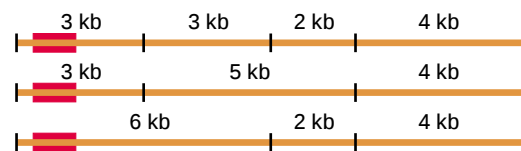
ment.) In the accompanying gel diagram, indicate the genotypes across the top and the phenotype (band position or positions) expected for each genotype. (The scale on the right shows the expected positions of fragments from 1 to 12 kb.)



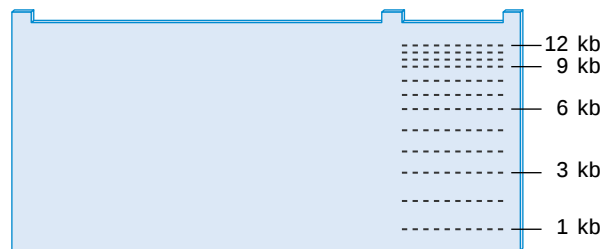
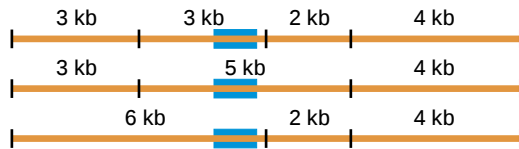
2.21 The accompanying diagram shows the DNA fragments associated with an RFLP revealed by a probe that hybridizes where shown by the rectangle. The tick marks are cleavage sites for the restriction enzyme used in the RFLP analysis. How many alleles does this RFLP have? How many genotypes are possible? In the accompanying gel diagram, indicate the phenotype (pattern of bands) expected of each genotype. (The scale on the right shows the expected positions of fragments from 1 to 12 kb.)



2.22 The thick horizontal lines shown below represent alternative DNA molecules at a particular locus in a human chromosome. The tick marks indicate the positions of restriction sites for a particular restriction enzyme. Genomic DNA from a sample of people is digested and analyzed by a Southern blot using a probe DNA that hybridizes at the position shown by the rectangle. How many possible RFLP alleles would be observed in the sample? How many genotypes?



2.23 The RFLPs described in Problem 2.22 are analyzed with the same restriction enzyme but a different probe, which hybridizes at the site indicated here by the rectangle. How many RFLP alleles would be found? How many genotypes? (Use the symbols A_1 , A_2 , ... to indicate the alleles.) In the accompanying gel diagram, indicate the genotypes across the top and the phenotype (band position or positions) expected for each genotype. (The scale on the right shows the expected positions of fragments from 1 to 12 kb.)



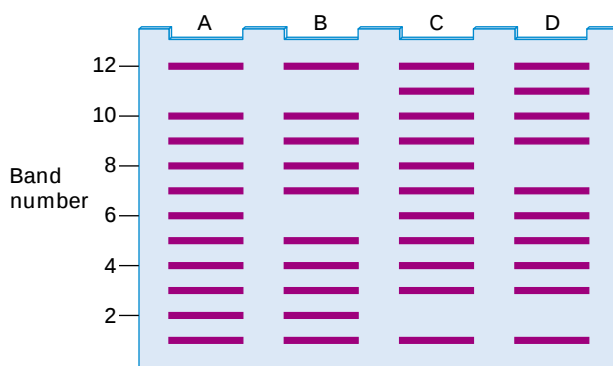
2.24 If hexamers were long enough oligonucleotides to serve as specific primers for PCR (for most purposes they are too short), what DNA fragment would be amplified using the “forward” primer pair 5'-AATGCC-3' and the “reverse” primer 3'-GCATGT-5' on the double-stranded DNA molecule shown below?

2.25 Would the primer pairs 3'-AATGCC-5' and 5'-GCATGT-3' amplify the same fragment described in the previous problem? Explain your answer.

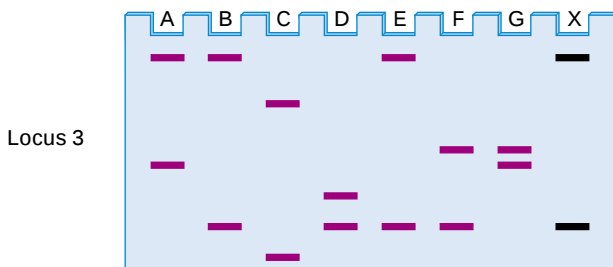
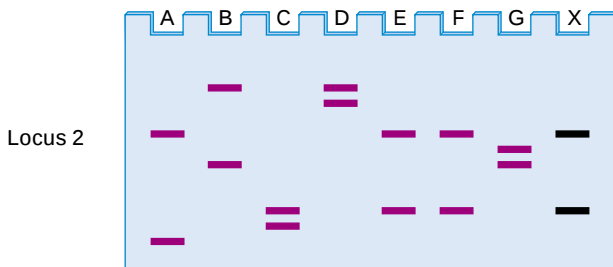
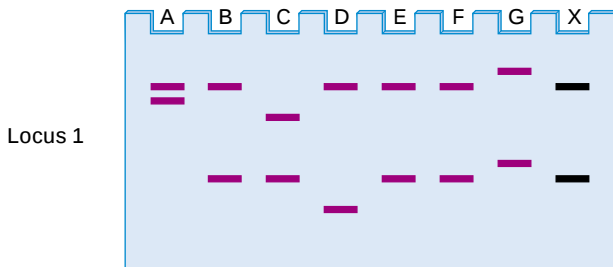
2.26 A human DNA fragment of 3 kb is to be amplified by PCR. The total genome size is 3×10^9 bp.

- Prior to amplification, what fraction of the total DNA does the target sequence constitute?
- What fraction does it constitute after 10 cycles of PCR?
- After 20 cycles of PCR?
- After 30 cycles of PCR?

2.27 RAPD analysis is carried out using genomic DNA from four individuals (A–D) sampled from a natural population of Hawaiian crickets. The gel shown at the top of the page resulted from PCR with one of the primer pairs tested. Which bands are the RAPD polymorphisms?



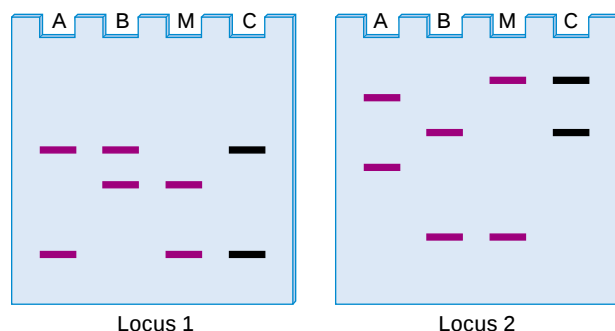
2.28 A cigarette butt found at the scene of a robbery is found to have a sufficient number of epithelial cells stuck to the paper for the DNA to be extracted and typed. Shown below are the results of typing for three probes (locus 1, locus 2, and locus 3) of the evidence (X) and 7 suspects (A through G). Which of the suspects can be excluded? Which cannot be excluded? Can you identify the robber? Explain your reasoning.



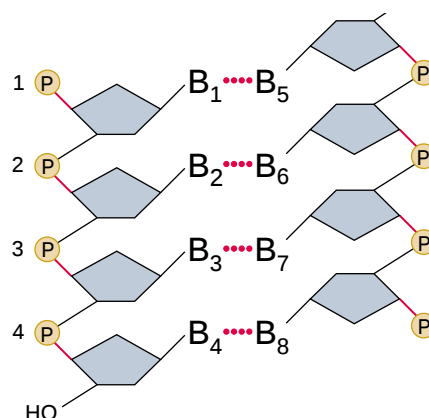
5'-GATTACCGGTAAATGCCGGATTAACCCGGTTATCAGGCCACGTACAACCTGGAGTCC-3'
3'-CTAATGGCCATTTACGGCCTAATTGGGCCAATAGTCCGGTGCATGTTGACCTCAGG-5'

Problem 2.24

2.29 A woman is uncertain which of two men is the father of her child. DNA typing is carried out on blood from the child (C), the mother (M), and each of the two males (A and B), using probes for a highly polymorphic DNA marker on two different chromosomes (“locus 1” and “locus 2”). The result is shown in the accompanying diagram. Can either male be excluded as the possible father? Explain your reasoning.



2.30 Snake venom phosphodiesterase cleaves the chemical bonds shown in red in the accompanying diagram, leaving mononucleotides that are phosphorylated in the 3' position. If the phosphates numbered 2 and 4 are radioactive, which mononucleotides will be radioactive after cleavage with snake venom phosphodiesterase?



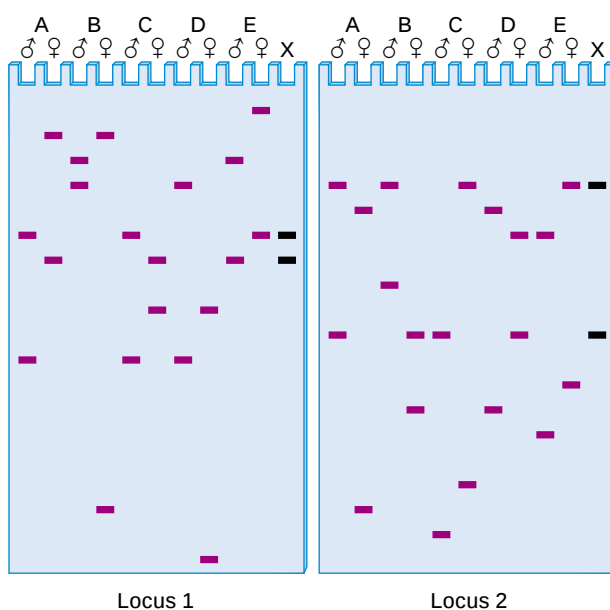
AU/ED: Art for prob 2.8 scaled to 85% for consistency. OK?

Challenge Problems

2.31 The genome of *Drosophila melanogaster* is 180×10^6 bp, and a fragment of size 1.8 kb is to be amplified by PCR. How many cycles of PCR are necessary for the amplified target sequence to constitute at least 99 percent of the total DNA?

2.32 A murder victim is found in an advanced state of decomposition and cannot be identified. Police suspect that the victim is one of five persons reported by their parents as missing. DNA typing is carried out on tissues from the victim (X) and on the five sets of parents (A through E), using probes for a highly polymorphic DNA marker on two different chromosomes (“locus 1” and “locus 2”). The result is shown in the diagram at the right. How do you interpret the fact that genomic DNA from each individual yields two bands? Can you identify the parents of the victim? Explain your reasoning.

2.33 The snake venom phosphodiesterase enzyme described in Problem 2.30 was originally used in a procedure called “nearest neighbor” analysis. In this procedure, a DNA strand is synthesized in the presence of all four trinucleotides, one of which carries a radioactive phosphate in the α (innermost) position. Then the DNA is digested to completion with snake venom phosphodiesterase, and the resulting mononucleotides are separated and assayed for radioactivity. Examine the diagram in Problem 2.30, and then answer the following questions.



Problem 2.32

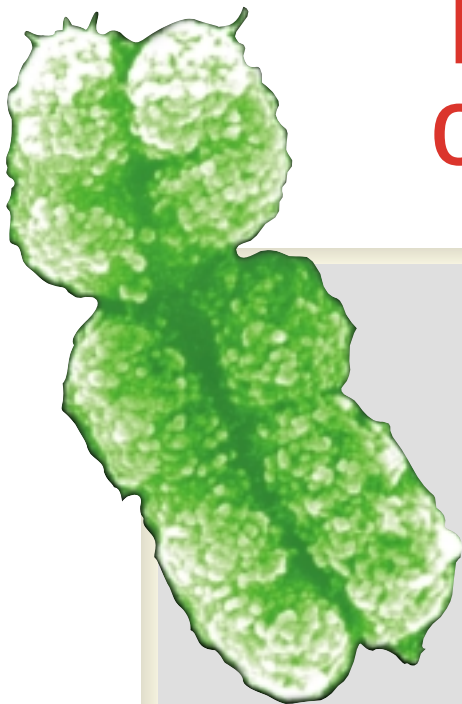
- How does this procedure reveal the “nearest neighbors” of the radioactive nucleotide?
- Is the “nearest neighbor” on the 5' or the 3' side of the labeled nucleotide?

Further Reading

- Botstein, D., R. L. White, M. Skolnick, and R. W. Davis. 1980. Construction of a genetic linkage map in man using restriction fragment length polymorphisms. *American Journal of Human Genetics* 32: 314.
- Calladine, C. R., and H. Drew. 1997. *Understanding DNA: The Molecule and How it Works*. 2d ed. San Diego: Academic Press.
- Cruzan, M. B. 1998. Genetic markers in plant evolutionary ecology. *Ecology* 79: 400.
- Danna, K., and D. Nathans. 1971. Specific cleavage of Simian Virus 40 DNA by restriction endonuclease of *Hemophilus influenzae*. *Proceedings of the National Academy of Sciences, USA* 68: 2913.
- DePamphilis, M. L., ed. 1996. *DNA Replication in Eukaryotic Cells*. Cold Spring Harbor, NY: Cold Spring Harbor Press.
- Eeles, R. A. and A. C. Stamps. 1993. *Polymerase Chain Reaction (PCR): The Technique and Its Application*. Austin TX: R. G. Landes.
- Frank-Kamenetskii, M. D. 1997. *Unraveling DNA: The Most Important Molecule of Life*. Tr. by L. Liapin. Reading MA: Addison-Wesley.
- Hartl, D. L. 2000. *A Primer of Population Genetics*. 3d ed. Sunderland, MA: Sinauer.
- Jorde, L. B., M. Bamshad, and A. R. Rogers. 1998. Using mitochondrial and nuclear DNA markers to reconstruct human evolution. *Bioessays* 20: 126.
- Kumar, L. S. 1999. DNA markers in plant improvement: An overview. *Biotechnology Advances* 17: 143.
- Loxdale, H. D., and G. Lushai. 1998. Molecular markers in entomology. *Bulletin of Entomological Research* 88: 577–600.
- Mitton, J. B. 1994. Molecular approaches to population biology. *Annual Review of Ecology & Systematics* 25: 45.
- Mullis, K. B. 1990. The unusual origin of the polymerase chain reaction. *Scientific American*, April.
- Olby, R. C. 1994. *The Path to the Double Helix: The Discovery of DNA*. New York: Dover.
- Pena, S. D. J., V. F. Prado, and J. T. Epplen. 1995. DNA diagnosis of human genetic individuality. *Journal of Molecular Medicine* 73: 555.
- Sayre, A. 1975. *Rosalind Franklin and DNA*. New York: Norton.
- Schafer, A. J., and J. R. Hawkins. 1998. DNA variation and the future of human genetics. *Nature Biotechnology* 16: 33.
- Smithies, O. 1995. Early days of electrophoresis. *Genetics* 139: 1.
- Thomson, G., and M. S. Esposito. 1999. The genetics of complex diseases. *Trends in Genetics* 15: M17.
- Wang, D. G., J.-B. Fan, C.-J. Siao, A. Berno, P. Young, R. Sapolsky *et al.* 1998. Large-scale identification, mapping, and genotyping of single-nucleotide polymorphisms in the human genome. *Science* 280: 1077.
- Watson, J. D. 1968. *The Double Helix*. New York: Atheneum.

CHAPTER 3

Transmission Genetics: The Principle of Segregation



CHAPTER OUTLINE

- 3.1 [Morphological and Molecular Phenotypes](#)
- 3.2 [Segregation of a Single Gene](#)
Phenotypic Ratios in the F₂ Generation
The Principle of Segregation
Verification of Segregation
The Testcross and the Backcross
- 3.3 [Segregation of Two or More Genes](#)
The Principle of Independent Assortment
The Testcross with Unlinked Genes
The Big Experiment
- 3.4 [Human Pedigree Analysis](#)
Characteristics of Dominant and Recessive Inheritance
Molecular Markers in Human Pedigrees
- 3.5 [Pedigrees and Probability](#)
Mutually Exclusive Possibilities
Independent Possibilities
- 3.6 [Incomplete Dominance and Epistasis](#)
Multiple Alleles
Human ABO Blood Groups
Epistasis
- 3.7 [Genetic Analysis: Mutant Screens and the Complementation Test](#)
The Complementation Test in Gene Identification
Why Does the Complementation Test Work?

PRINCIPLES

- Inherited traits are determined by the genes present in the reproductive cells united in fertilization.
- Genes are usually inherited in pairs—one from the mother and one from the father.
- The genes in a pair may differ in DNA sequence and in their effect on the expression of a particular inherited trait.
- The maternally and paternally inherited genes are not changed by being together in the same organism.
- In the formation of reproductive cells, the paired genes separate again into different cells.
- Random combinations of reproductive cells containing different genes result in Mendel's ratios of traits appearing among the progeny.
- The ratios actually observed for any trait are determined by the types of dominance and gene interaction.
- In genetic analysis, the complementation test is used to determine whether two recessive mutations that cause a similar phenotype are alleles of the same gene. The mutant parents are crossed, and the phenotype of the progeny is examined. If the progeny phenotype is nonmutant (complementation occurs), the mutations are in different genes; if the progeny phenotype is mutant (complementation does not occur), the mutations are in the same gene.

CONNECTIONS

What Did Gregor Mendel Think He Discovered?

Gregor Mendel 1866
Experiments on Plant Hybrids

Troubadour

The Huntington's Disease Collaborative Research Group 1993
A Novel Gene Containing a Trinucleotide Repeat That Is Expanded and Unstable on Huntington's Disease Chromosomes

AU: Reference in first paragraph changed to Fig. 2.25. Correct? or restore to Fig. 2.24?

AU: In first paragraph, phrase “phenotype is characterized by” changed to “phenotype consists.” OK?

In this small garden plot adjacent to the monastery of St. Thomas, Gregor Mendel grew more than 33,500 pea plants in the years 1856–1863, including more than 6400 plants in one year alone. He received some help from two fellow monks who assisted in the experiments.



In Chapter 2 we emphasized that in modern genetics, a typical phenotype consists of a band present at a particular position in a DNA electrophoresis gel. In a method of genetic analysis such as RAPD (random amplified polymorphic DNA), genomic DNA from a single individual can yield 30 or more bands (see Figure 2.25). Each of these bands represents a phenotype and results from the PCR (polymerase chain reaction) amplification of a single region of DNA in the genome of the individual.

In this chapter we consider how genes are transmitted from parents to offspring and how this determines the distribution of genotypes and phenotypes among related individuals. The study of the inheritance of traits constitutes **transmission genetics**. This subject is also called **Mendelian genetics** because the underlying principles were first deduced from experiments in garden peas (*Pisum sativum*) carried out in the years 1856–1863 by Gregor Mendel, a monk at the monastery of St. Thomas in the town of Brno (Brünn), in the Czech Republic. He reported his experiments to a local natural history society, published the results and his interpretation in its scientific journal in 1866, and began exchanging letters with one of the leading botanists of the time. His experiments were careful and exceptionally well documented, and his paper contains the first clear exposition of the statistical rules governing the transmission of genes from generation to generation. Nevertheless, Mendel's paper was ignored for 34 years until its significance was finally appreciated.

3.1 Morphological and Molecular Phenotypes

Geneticists study traits in which one organism differs from another in order to discover the genetic basis of the difference. Until the advent of molecular genetics, geneticists dealt mainly with morphological traits, in which the differences between organisms can be expressed in terms of color, shape, or size. Mendel studied seven morphological traits contrasting in seed shape, seed color, flower color, pod shape, and so forth (Figure 3.1). Perhaps the most widely known example of a contrasting Mendelian trait is round versus wrinkled seeds. As pea seeds dry, they lose water and shrink. Round seeds are round because they shrink uniformly; wrinkled seeds are wrinkled because they shrink irregularly. The wrinkled phenotype is due to the absence of a branched-chain form of starch known as *amylopectin*, which is not synthesized in wrinkled seeds owing to a defect in the enzyme starch-branching enzyme I (SBEI).

The nonmutant, or **wildtype**, allele of the gene for SBEI is designated *W* and the mutant allele *w*. Seeds that are heterozygous *Ww* have only half as much SBEI enzyme as wildtype homozygous *WW* seeds, but this half the normal amount of enzyme produces enough amylopectin for the heterozygous *Ww* seeds to shrink uniformly and remain phenotypically round. Hence, with respect to seed shape, the wildtype *W* allele is dominant over the mutant *w* allele. Because the *w* allele is recessive, only the homozygous *ww* seeds become wrinkled.

The molecular basis of the wrinkled mutation is that the SBEI gene has become interrupted by the insertion, into the gene, of a DNA sequence called a **transposable element**. Such DNA sequences are capable of moving (*transposition*) from one location to another within a chromosome or between chromosomes. The molecular mechanism of transposition is discussed in Chapter 7, but for our present purposes it is necessary to know only that transposable elements are present in most genomes, especially the large genomes of eukaryotes, and that many spontaneous mutations result from the insertion of transposable elements into a gene.

	Parental strain 1: Dominant	Parental strain 2: Recessive	Phenotype of progeny of monohybrid cross
Seed shape	 Round	 Wrinkled	 Round
Seed color	 Yellow	 Green	 Yellow
Flower color	 Purple	 White	 Purple
Pod shape	 Inflated	 Constricted	 Inflated
Pod color	 Green	 Yellow	 Green
Flower and pod position	 Axial (along stem)	 Terminal (at top of stem)	 Axial
Stem length	 Standard	 Dwarf	 Standard

Figure 3.1 The seven different traits studied in peas by Mendel. The phenotype shown at the far right is the dominant trait, which appears in the hybrid produced by crossing.

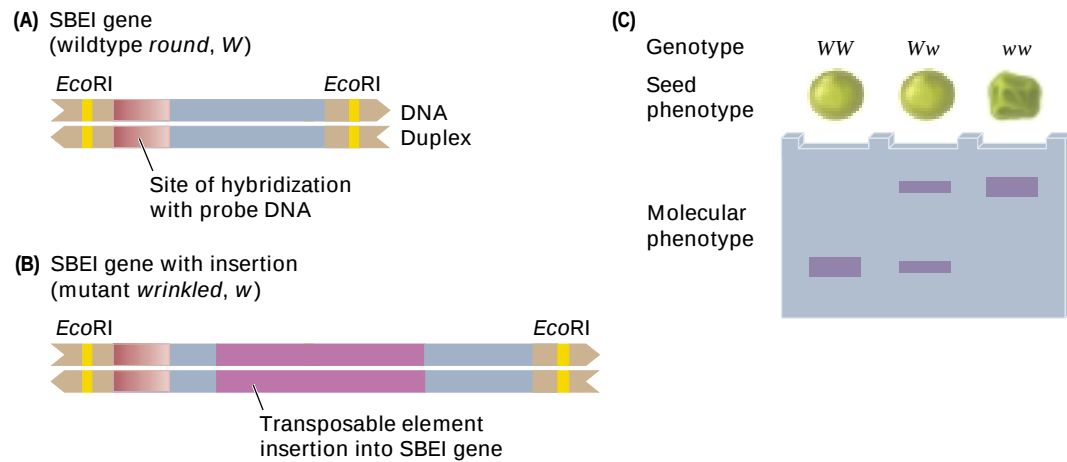


Figure 3.2 (A) *W* (round) is an allele of a gene that specifies the amino acid sequence of starch branching enzyme I (SBEI). (B) *w* (wrinkled) is an allele that encodes an inactive form of the enzyme because its DNA sequence is interrupted by the insertion of a transposable element. (C) At the level of the morphological phenotype, *W* is dominant to *w*: Genotype *WW* and *Ww* have round seeds, whereas genotype *ww* has wrinkled seeds. The molecular difference between the alleles can be detected as a restriction fragment length polymorphism (RFLP) using the enzyme *EcoRI* and a probe that hybridizes at the site shown. At the molecular level, the alleles are codominant: DNA from each genotype yields a different molecular phenotype—a single band differing in size for homozygous *WW* and *ww*, and both bands for heterozygous *Ww*.

Figure 3.2 includes a diagram of the DNA structure of the wildtype *W* and mutant *w* alleles and shows the DNA insertion that interrupts the *w* allele. Highlighted are two *EcoRI* restriction sites, present in both alleles, that flank the site of the insertion. The diagram in part C indicates what pattern of bands would be expected if one were to carry out a Southern blot (Section 2.3) in which genomic DNA was digested to completion with *EcoRI* and then the resulting fragments were separated by electrophoresis and hybridized with a labeled probe complementary to a region shared between the *W* and *w* alleles. The *EcoRI* fragment from the *W* allele would be smaller than that of the *w* allele, because of the inserted DNA in the *w* allele, and thus it would migrate faster than the corresponding fragment from the *w* allele and move to a position closer to the bottom of the gel. Genomic DNA from homozygous *WW* would yield a single, faster-migrating band; that from homozygous *ww* a single, slower-migrating band; and that from heterozygous *Ww* two bands with the same electrophoretic mobility as those observed in the homozygous genotypes. In Figure 3.2C, the band in the homozygous genotypes is shown as somewhat thicker than those from the heterozygous genotype, because the single band in each ho-

mozygous genotype comes from the two copies of whichever allele is homozygous, and thus it contains more DNA than the corresponding DNA in the heterozygous genotype, in which only one copy of each allele is present.

Hence, as illustrated in Figure 3.2C, the RFLP analysis clearly distinguishes between the genotypes *WW*, *Ww*, and *ww*, because the heterozygous *Ww* genotype exhibits both of the bands observed in the homozygous genotypes. This situation is described by saying that the *W* and *w* alleles are **codominant** with respect to the molecular phenotype. However, as indicated by the seed shapes in Figure 3.2C, *W* is dominant over *w* with respect to morphological phenotype. In the discussion that follows, we use the RFLP analysis of the *W* and *w* alleles to emphasize the importance of molecular phenotypes in modern genetics and to demonstrate experimentally the principles of genetic transmission.

3.2 Segregation of a Single Gene

Mendel selected peas for his experiments for two primary reasons. First, he had access to varieties that differed in contrasting traits, such as round versus wrinkled seeds

and yellow versus green seeds. Second, his earlier studies had indicated that peas usually reproduce by self-pollination, in which pollen produced in a flower is used to fertilize the eggs in the same flower. To produce hybrids by cross-pollination (**outcrossing**), all he needed to do was open the keel petal (enclosing the reproductive structures), remove the immature anthers (the pollen-producing structures) before they shed pollen, and dust the stigma (the female structure) with pollen taken from a flower on another plant (Figure 3.3). The relatively small space needed to grow each plant, and the relatively large number of progeny that

could be obtained, gave him the opportunity, as he says in his paper, to “determine the number of different forms in which hybrid progeny appear” and to “ascertain their numerical interrelationships.”

The fact that garden peas are normally self-fertilizing means that in the absence of deliberate outcrossing, plants with contrasting traits are usually homozygous for alternative alleles of a gene affecting the trait; for example, plants with round seeds have genotype WW and those with wrinkled seeds genotype ww . The homozygous genotypes are indicated experimentally by the observation that hereditary traits in each

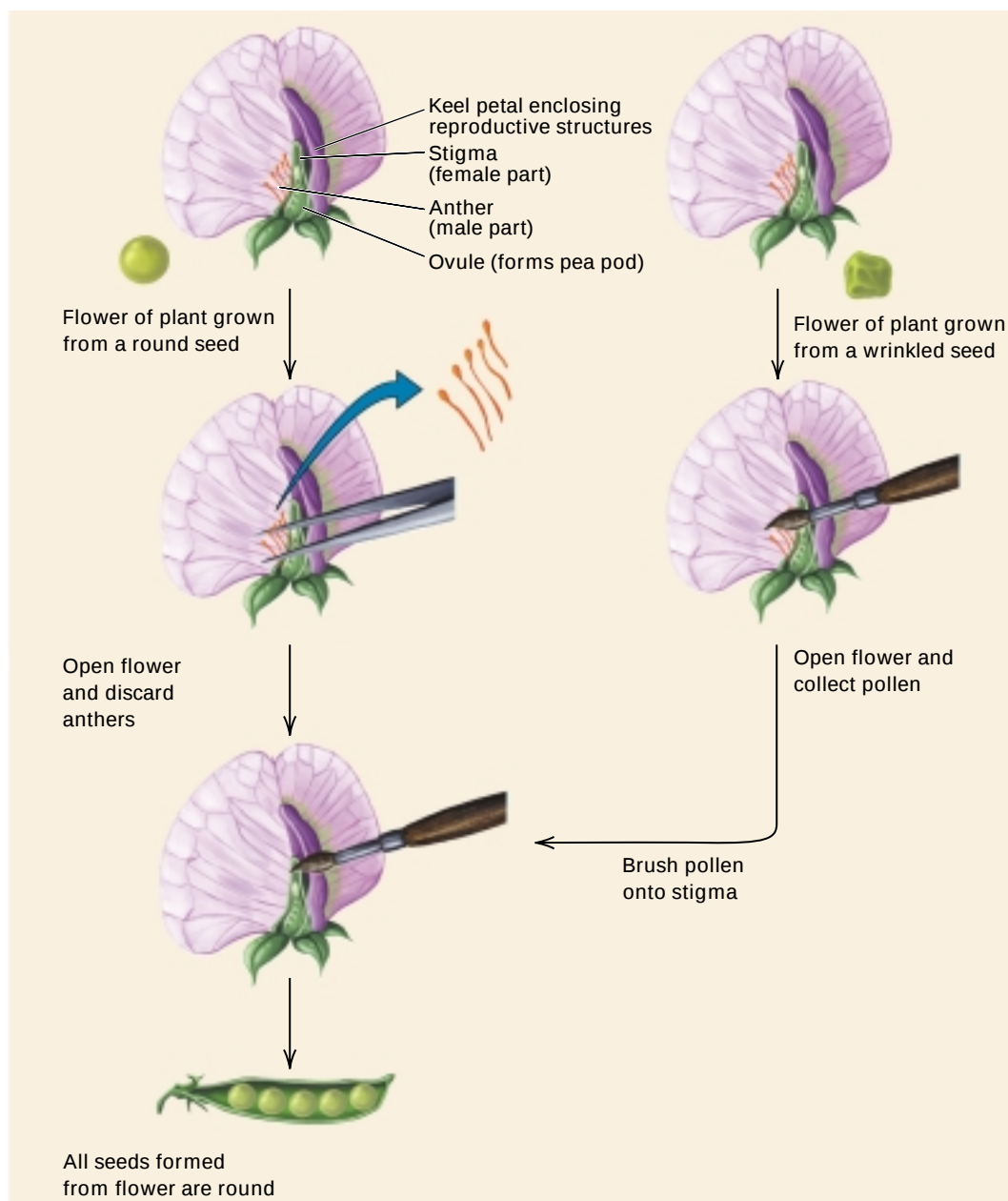


Figure 3.3 Crossing pea plants requires some minor surgery in which the anthers of a flower are removed before they produce pollen. The stigma, or female part of the flower, is not removed. It is fertilized by brushing with mature pollen grains taken from another plant.

variety are **true-breeding**, which means that plants produce only progeny like themselves when allowed to self-pollinate normally. Outcrossing between plants that differ in one or more traits creates a **hybrid**. If the parents differ in one, two, or three traits, the hybrid is a *monohybrid*, *dihybrid*, or *trihybrid*. In keeping track of parents and their hybrid progeny, we say that the parents constitute the **P₁ generation** and their hybrid progeny the **F₁ generation**.

Because garden peas are sexual organisms, each cross can be performed in two ways, depending on which phenotype is present in the female parent and which in the male parent. With round versus wrinkled, for example, the female parent can be round and the male wrinkled, or the reverse. These are called **reciprocal crosses**. Mendel was the first to demonstrate the following important principle:

The outcome of a genetic cross does not depend on which trait is present in the male and which is present in the female; reciprocal crosses yield the same result.

This principle is illustrated for round and wrinkled seeds in Figure 3.4. The gel icons show the RFLP bands that genomic DNA from each type of seed in these crosses would yield. The genotypes of the crosses and their progeny are as follows:

Cross A: $WW / \times ww ? \rightarrow Ww$

Cross B: $ww / \times WW ? \rightarrow Ww$

In both reciprocal crosses, the progeny have the morphological phenotype of round seeds, but as shown by the RFLP diagrams on the right, all progeny genotypes are actually heterozygous Ww and therefore different from either parent. The genetic equivalence of reciprocal crosses illustrated in Figure 3.4 is a principle that is quite general in its applicability, but there is an important exception, having to do with sex chromosomes, that will be discussed in Chapter 4.

In the following section we examine a few of Mendel's original experiments in the context of RFLP analysis in order to relate the morphological phenotypes and their

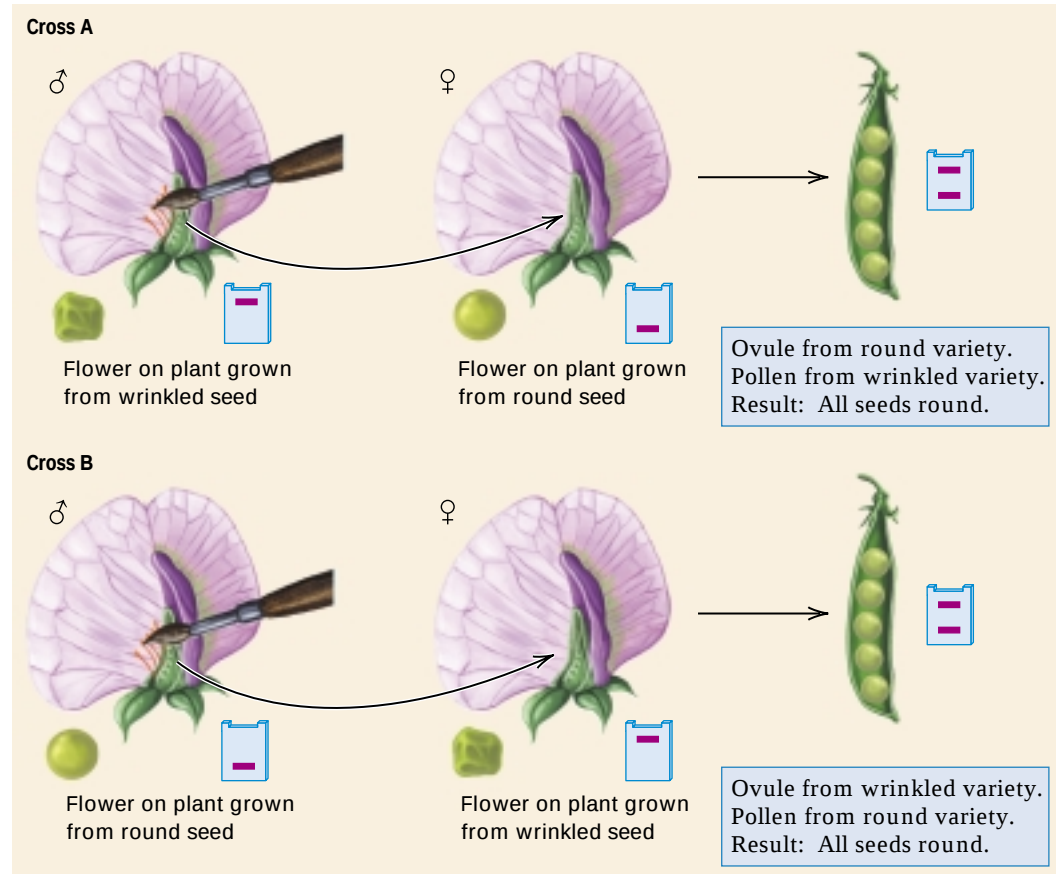


Figure 3.4 Morphological and molecular phenotypes showing the equivalence of reciprocal crosses. In this example, the hybrid seeds are round and yield an RFLP pattern with two bands, irrespective of the direction of the cross.

Phenotypic Ratios in the F₂ Generation

The diagram illustrates Mendel's experiment on the inheritance of seed shape in pea plants. It is divided into three horizontal sections representing different generations:

- Parental lines:** This section shows two pea plants. On the left is a **Round** plant, and on the right is a **Wrinkled** plant. An arrow points from the space between them down to the F₁ generation.
- F₁ generation:** This section shows the result of the cross: **All round seeds**. This indicates that the round trait is dominant.
- F₂ generation:** This section shows the results of self-crossing the F₁ plants. It displays a **3 round (5474) : 1 wrinkled (1850)** phenotypic ratio, which corresponds to a 3:1 genotypic ratio.

A 3 : 1 ratio of dominant : recessive forms in the F₂ progeny is characteristic of simple Mendelian inheritance. Mendel's data demonstrating this point are shown in [Table 3.1](#). Note that the first two traits in the table (round versus wrinkled seeds and yellow versus green seeds) have many more observations than any of the others. The reason is that these traits can be classified directly in the seeds, whereas the others can be classified only in the mature plants, and Mendel could analyze many more seeds than he could mature plants. The

Parental traits	F ₁ trait	Number of F ₂ progeny	F ₂ ratio
Round × wrinkled (seeds)	Round	5474 round, 1850 wrinkled	2.96 : 1
Yellow × green (seeds)	Yellow	6022 yellow, 2001 green	3.01 : 1
Purple × white (flowers)	Purple	705 purple, 224 white	3.15 : 1
Inflated × constricted (pods)	Inflated	882 inflated, 299 constricted	2.95 : 1
Green × yellow (unripe pods)	Green	428 green, 152 yellow	2.82 : 1
Axial × terminal (flower position)	Axial	651 axial, 207 terminal	3.14 : 1
Long × short (stems)	Long	787 long, 277 short	2.84 : 1

- The F_1 hybrids express only the dominant trait (because the F_1 progeny are heterozygous—for example, Ww)
- In the F_2 generation, plants with either the dominant or the recessive trait are present (which means that some F_2 progeny are homozygous—for example, ww).
- In the F_2 generation, there are approximately three times as many plants with the dominant phenotype as plants with the recessive phenotype.

The Principle of Segregation

The 3 : 1 ratio can be explained with reference to **Figure 3.5**. This is the heart of Mendelian genetics. You should master it and be able to use it to deduce the progeny types produced in crosses. The diagram illustrates these key features of single-gene inheritance:

1. Genes come in pairs, which means that a cell or individual has *two* copies (alleles) of each gene.
2. For each pair of genes, the alleles may be identical (homozygous WW or homozygous ww), or they may be different (heterozygous Ww).

3. Each reproductive cell (**gamete**) produced by an individual contains only *one* allele of each gene (that is, either W or w).
4. In the formation of gametes, any particular gamete is equally likely to include either allele (hence, from a heterozygous Ww genotype, half the gametes contain W and the other half contain w).
5. The union of male and female reproductive cells is a random process that reunites the alleles in pairs.

The essential feature of transmission genetics is the separation, technically called **segregation**, in unaltered form, of the two alleles in an individual during the formation of its reproductive cells. Segregation

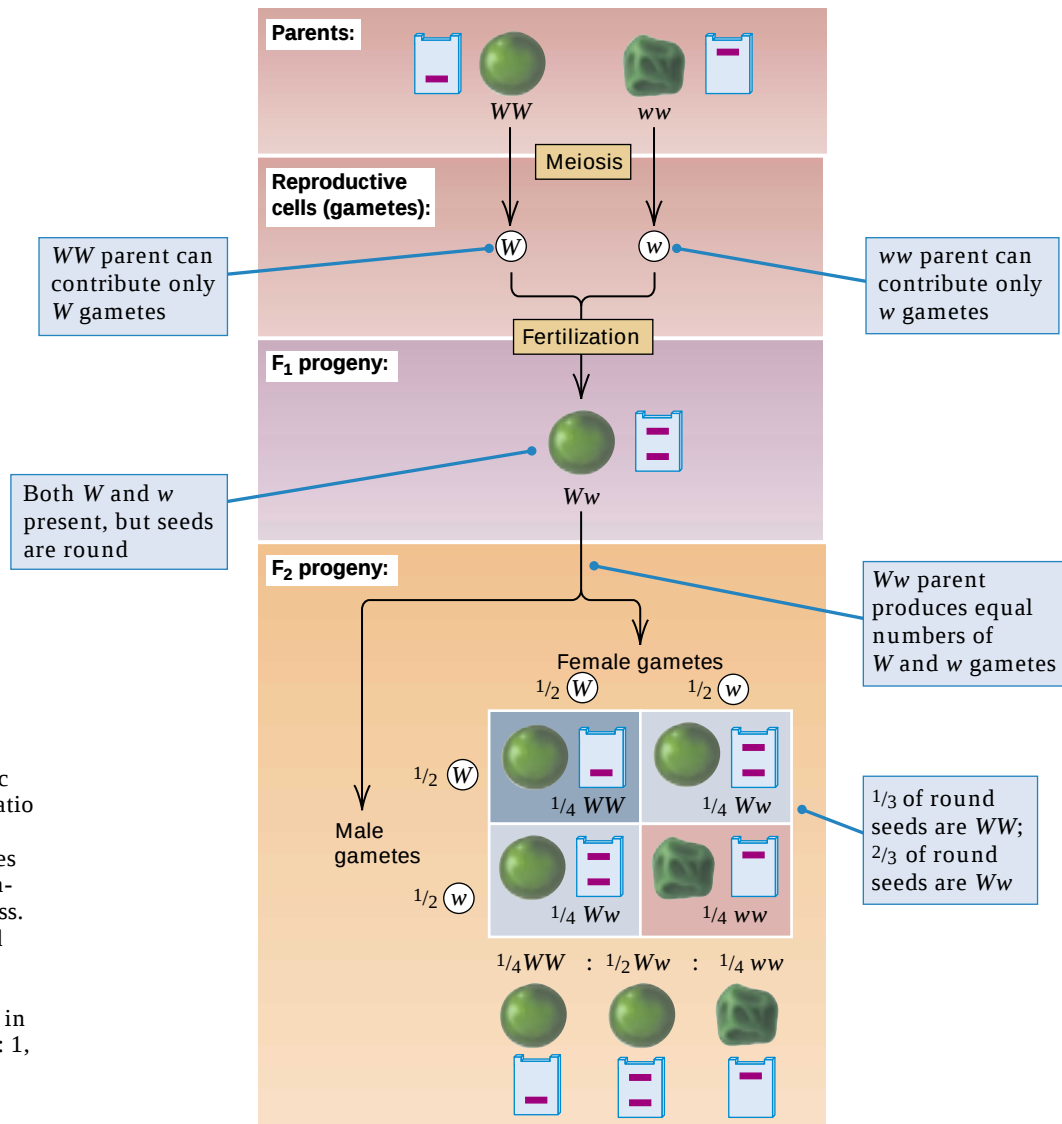


Figure 3.5 A diagrammatic explanation of the 3 : 1 ratio of dominant : recessive morphological phenotypes observed in the F_2 generation of a monohybrid cross. The 3 : 1 ratio is observed because of dominance. Note that the ratio of $WW : Ww : ww$ genotypes in the F_2 generation is 1 : 2 : 1, as can be seen from the restriction fragment phenotypes.

corresponds to points 3 and 4 in the foregoing list. The principle of segregation is sometimes called *Mendel's first law*.

The Principle of Segregation: In the formation of gametes, the paired hereditary determinants separate (segregate) in such a way that each gamete is equally likely to contain either member of the pair.

Another key feature of transmission genetics is that the hereditary determinants are present as pairs in both the parental organisms and the progeny organisms but as single copies in the reproductive cells. This feature corresponds to points 1 and 5 in the foregoing list.

Figure 3.5 illustrates the biological mechanism underlying the important Mendelian ratios in the F_2 generation of 3 : 1 for phenotypes and 1 : 2 : 1 for genotypes. To understand these ratios, consider first the parental generation in which the original cross is $WW \times ww$. The sex of the parents is not stated because reciprocal crosses yield identical results. (There is, however, a convention in genetics that unless otherwise specified, crosses are given with the female parent listed first.) In the original cross, the WW parent produces only W -containing gametes, whereas the ww parent produces only w -containing gametes. Segregation still takes place in the homozygous genotypes as well as in the heterozygous genotype, even though all of the gametes carry the same type of allele (W from homozygous WW and w from homozygous ww). When the W -bearing and w -bearing gametes come together in fertilization, the hybrid genotype is heterozygous Ww , which is shown by the bands in the gel icon next to the F_1 progeny. With regard to seed shape, the hybrid Ww seeds are round because W is dominant over w .

When the heterozygous F_1 progeny form gametes, segregation implies that half the gametes will contain the W allele and the other half will contain the w allele. These gametes come together at random when an F_1 individual is self-fertilized or when two F_1 individuals are crossed. The result of random fertilization can be deduced from the sort of cross-multiplication square shown at the bottom of the figure, in which the female gametes and their fre-

quencies are arrayed across the top margin and the male gametes and their frequencies along the left-hand margin. This calculating device is widely used in genetics and is called a **Punnett square** after its inventor Reginald C. Punnett (1875–1967). The Punnett square in Figure 3.5 shows that random combinations of the F_1 gametes result in an F_2 generation with the genotypic composition $1/4 WW$, $1/2 Ww$, and $1/4 ww$. This can be confirmed by the RFLP banding patterns in the gel icons because of the codominance of W and w with respect to the molecular phenotype. But because W is dominant over w with respect to the morphological phenotype, the WW and Ww genotypes have round seeds and the ww genotypes have wrinkled seeds, yielding the phenotypic ratio of round : wrinkled seeds of 3 : 1. Hence it is a combination of segregation, random union of gametes, and dominance that results in the 3 : 1 ratio.

The ratio of F_2 genotypes is as important as the ratio of F_2 phenotypes. The Punnett square in Figure 3.5 also shows that the ratio of $WW : Ww : ww$ genotypes is 1 : 2 : 1, which can be confirmed directly by the RFLP analysis.

Verification of Segregation

The round seeds in Figure 3.5 conceal a genotypic ratio of 1 WW : 2 Ww . To say the same thing in another way, among the F_2 seeds that are round (or, more generally, among organisms that show the dominant morphological phenotype), $1/3$ are homozygous (in this example, WW) and $2/3$ are heterozygous (in this example, Ww). This conclusion is obvious from the RFLP patterns in Figure 3.5, but it is not at all obvious from the morphological phenotypes. Unless you knew something about genetics already, it would be a very bold hypothesis, because it implies that two organisms with the same morphological phenotype (in this case round seeds) might nevertheless differ in molecular phenotype and in genotype.

Yet this is exactly what Mendel proposed. But how could this hypothesis be tested experimentally? He realized that it could be tested via self-fertilization of the F_2 plants. With self-fertilization, plants grown

What Did Gregor Mendel Think He Discovered?

Gregor Mendel 1866

Monastery of St. Thomas, Brno
[then Brünn], Czech Republic
Experiments on Plant Hybrids
(original in German)

Mendel's paper is remarkable for its precision and clarity. It is worth reading in its entirety for this reason alone. Although the most important discovery attributed to Mendel is segregation, he never uses this term. His description of segregation is found in the first passage in italics in the excerpt. (All of the italics are reproduced from the original.) In his description of the process, he takes us carefully through the separation of A and a in gametes and their coming together again at random in fertilization. One flaw in the description is Mendel's occasional confusion between genotype and phenotype, which is illustrated by his writing A instead of AA and a instead of aa in the display toward the end of the passage. Most early geneticists

made no consistent distinction between genotype and phenotype until 1909, when the terms themselves were coined.

Artificial fertilization undertaken on ornamental plants to obtain new color variants initiated the experiments reported here. The striking regularity with which the same hybrid forms always reappeared whenever fertilization between like species took place suggested further experiments whose task it was to follow that development of hybrids in their progeny. . . . This paper discusses the attempt at such a detailed experiment. . . . Whether the plan by which the individual experiments were set up and carried out was adequate to the assigned task should be decided by a benevolent judgment. . . . [Here the experimental results are described in detail.] Thus

experimentation also justifies the assumption that pea hybrids form germinal and pollen cells that in their composition correspond in equal numbers to all the constant forms resulting from the combination of traits united through fertilization.

Whether the plan by which the individual experiments were set up and carried out was adequate to the assigned task should be decided by a benevolent judgment.

The difference of forms among the progeny of hybrids, as well as the ratios in which they are observed, find an adequate explanation in the principle [of segregation] just deduced. The simplest case is given by the series for one pair of differing

traits. It is shown that this series is described by the expression: $A + 2Aa + a$, in which A and a signify the forms with constant differing traits, and Aa the form hybrid for both. The series contains four individuals in three different terms. In their production, pollen and germinal

from the homozygous WW genotypes should be true-breeding for round seeds, whereas those from the heterozygous Ww genotypes should yield round and wrinkled seeds in the ratio of 3 : 1. On the other hand, the plants grown from wrinkled seeds should be true-breeding for wrinkled because these plants are homozygous ww. The results Mendel obtained are summarized in Figure 3.6. As predicted from the genetic hypothesis, plants grown from F₂ wrinkled seeds were true-breeding for wrinkled seeds, yielding only wrinkled seeds in the F₃ generation. But some of the plants grown from round seeds showed evidence of segregation. Among 565 plants grown from round F₂ seeds, 372 plants produced both round and wrinkled seeds in a proportion very close to 3 : 1, whereas the remaining 193 plants produced only round seeds in

the F₃ generation. The ratio 193 : 372 equals 1 : 1.93, which is very close to the ratio 1 : 2 of WW : Ww genotypes predicted from the genetic hypothesis.

An important feature of the homozygous round and homozygous wrinkled seeds produced in the F₂ and F₃ generations is that the phenotypes are exactly the same as those observed in the original parents in the P₁ generation. This makes sense in terms of DNA, because the DNA of each allele remains unaltered unless a new mutation happens to occur. Mendel described this result in a letter by saying that in the progeny of crosses, “the two parental traits appear, separated and unchanged, and there is nothing to indicate that one of them has either inherited or taken over anything from the other.” From this finding, he concluded that the hereditary determinants for

cells of form *A* and *a* participate, on the average, equally in fertilization; therefore each form manifests itself twice, since four individuals are produced. Participating in fertilization are thus:

Pollen cells *A* + *A* + *a* + *a*
 Germinal cells *A* + *A* + *a* + *a*

It is entirely a matter of chance which of the two kinds of pollen combines with each single germinal cell. However, according to the laws of probability, in an average of many cases it will always happen that every pollen form *A* and *a* will unite equally often with every germinal-cell form *A* and *a*; therefore, in fertilization, one of the two pollen cells *A* will meet a germinal cell *A*, the other a germinal cell *a*, and equally, one pollen cell *a* will become associated with a germinal cell *A*, and the other *a*.

Pollen cells *A* *A* *a* *a*
 j *j* *j* *j*
 Germinal cells *A* *A* *a* *a*

The result of fertilization can be visualized by writing the designations for associated germinal and pollen cells in the form of fractions, pollen cells above the line, germinal cells below. In the case under discussion one obtains

$$\frac{A}{A} + \frac{A}{a} + \frac{a}{A} + \frac{a}{a}$$

In the first and fourth terms germinal and pollen cells are alike; therefore the products of their association must be constant, namely *A* and *a*; in the second and third, however, a union of the two differing parental traits takes place again, therefore the forms arising from such fertilizations are absolutely identical with the hybrid from which they derive. Thus, repeated hybridization takes place. The striking phenomenon, that hybrids are able to produce, in addition to the two parental types, progeny that resemble themselves is thus explained: *Aa* and *aA* both give the same association, *Aa*, since, as mentioned earlier, it

makes no difference to the consequence of fertilization which of the two traits belongs to the pollen and which to the germinal cell. Therefore

$$\frac{A}{A} + \frac{A}{a} + \frac{a}{A} + \frac{a}{a} = A + 2Aa + a$$

This represents the average course of self-fertilization of hybrids when two differing traits are associated in them. In individual flowers and individual plants, however, the ratio in which the members of the series are formed may be subject to not insignificant deviations. . . . Thus it was proven experimentally that, in *Pisum*, hybrids form different kinds of germinal and pollen cells and that this is the reason for the variability of their offspring.

Source: *Verhandlungen des naturforschenden Vereines in Brünn* 4: 3–47.

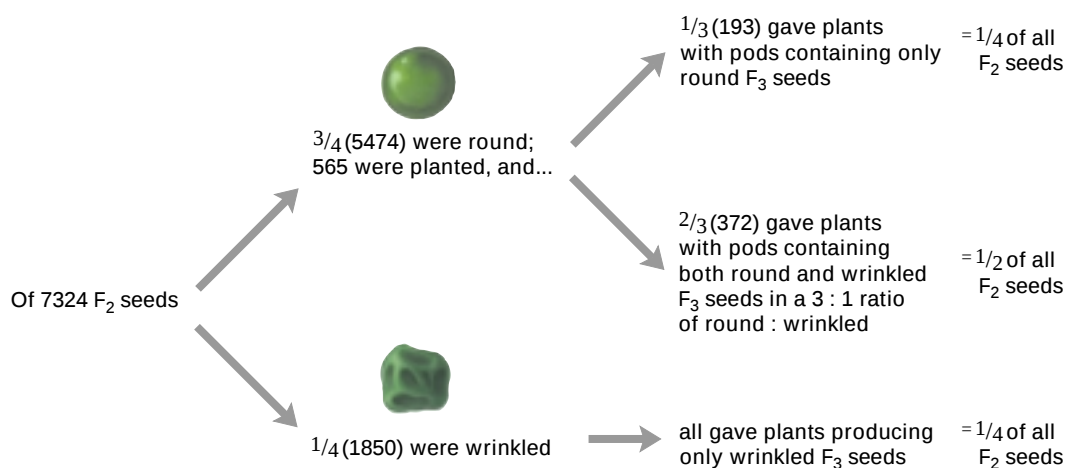


Figure 3.6 Summary of *F*₂ phenotypes and the progeny produced by self-fertilization.

the traits in the parental lines were transmitted as two different elements that retain their purity in the hybrids. In other words, the hereditary determinants do not “mix”

or “contaminate” each other. In modern terminology, this means that, with rare but important exceptions, genes are transmitted unchanged from generation to generation.

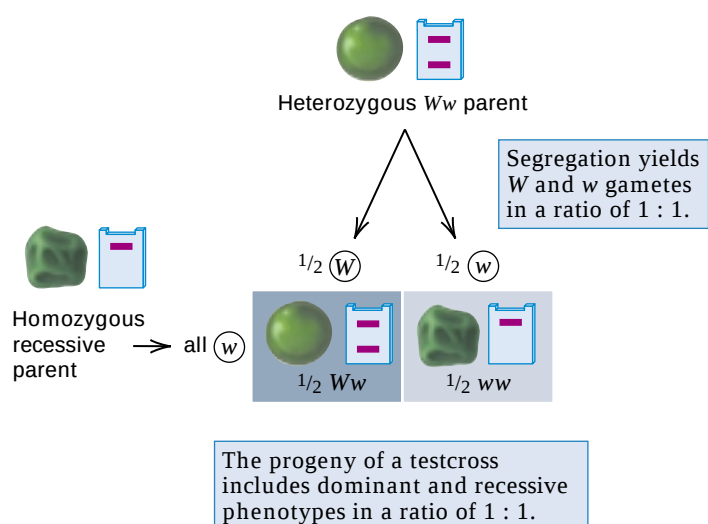


Figure 3.7 In a testcross of an Ww heterozygous parent with a ww homozygous recessive, the progeny are Ww and ww in the ratio of 1 : 1. A testcross shows the result of segregation.

The Testcross and the Backcross

Another straightforward way of testing the genetic hypothesis in Figure 3.5 is by means of a **testcross**, a cross between an organism that is heterozygous for one or more genes (for example, Ww) and an organism that is homozygous for the recessive alleles (for example, ww). The result of such a testcross is shown in Figure 3.7. Because the heterozygous parent is expected to produce W and w gametes in equal numbers, whereas the homozygous recessive produces only w gametes, the expected progeny are $\frac{1}{2}$ with the genotype Ww and $\frac{1}{2}$ with the genotype ww . The former have the dominant phenotype because W is dominant over w , and the latter have the recessive phenotype. A testcross is often extremely useful in genetic analysis:

In a testcross, the phenotypes of the progeny reveal the relative frequencies of the different

gametes produced by the heterozygous parent, because the recessive parent contributes only recessive alleles.

Mendel carried out a series of testcrosses with various traits. The results are shown in Table 3.2. In all cases, the ratio of phenotypes among the testcross progeny is very close to the 1 : 1 ratio expected from segregation of the alleles in the heterozygous parent.

Another valuable type of cross is a **backcross**, in which hybrid organisms are crossed with one of the parental genotypes. Backcrosses are commonly used by geneticists and by plant and animal breeders, as we will see in later chapters. Note that the testcrosses in Table 3.2 are also backcrosses, because in each case, the F_1 heterozygous parent came from a cross between the homozygous dominant and the homozygous recessive.

3.3 Segregation of Two or More Genes

The results of many genetic crosses depend on the segregation of the alleles of two or more genes. The genes may be in different chromosomes or in the same chromosome. Although in this section we consider the case of genes that are in two different chromosomes, the same principles apply to genes that are in the same chromosome but are so far distant from each other that they segregate independently. The case of *linkage* of genes in the same chromosome is examined in Chapter 5.

To illustrate the principles, we consider again a cross between homozygous genotypes, but in this case homozygous for the alleles of two genes. A specific example is a true-breeding variety of garden peas with seeds that are wrinkled and green (genotype $ww\ gg$) versus a variety with seeds that are round and yellow (genotype $WW\ GG$). As suggested by the use of uppercase and lowercase symbols for the alleles, the dominant alleles are W and G , the recessive alleles w and g . Crossing these strains yields F_1 seeds with the genotype $Ww\ Gg$, which are phenotypically round and yellow because of the dominance relations. When the F_1 seeds are grown into mature plants and self-fertilized, the F_2 progeny show the result of simultaneous segregation of the W , w allele pair and the G , g allele pair. When

Table 3.2 Mendel's testcross results

Testcross (F_1 heterozygote $\times$ homozygous recessive)	Progeny from testcross	Ratio
Round $\times$ wrinkled seeds	193 round, 192 wrinkled	1.01 : 1
Yellow $\times$ green seeds	196 yellow, 189 green	1.04 : 1
Purple $\times$ white flowers	85 purple, 81 white	1.05 : 1
Long $\times$ short stems	87 long, 79 short	1.10 : 1

Mendel performed this cross, he obtained the following numbers of F₂ seeds:

Round, yellow	315
Round, green	108
Wrinkled, yellow	101
Wrinkled, green	32
Total	556

In these data, the first thing to be noted is the expected 3 : 1 ratio for each trait considered separately. The ratio of round : wrinkled (pooling across yellow and green) is

$$\begin{aligned} & (315 + 108) : (101 + 32) \\ & = 423 : 133 \\ & = 3.18 : 1 \end{aligned}$$

And the ratio of yellow : green (pooling across round and wrinkled) is

$$\begin{aligned} & (315 + 101) : (108 + 32) \\ & = 416 : 140 \\ & = 2.97 : 1 \end{aligned}$$

Both of these ratios are in satisfactory agreement with 3 : 1. (Testing for goodness of fit to a predicted ratio is described in Chapter 4.) Furthermore, in the F₂ progeny of the dihybrid cross, the separate 3 : 1 ratios for the two traits were combined at random. With random combinations, as shown in Figure 3.8, among the 3/4 of the progeny that are round, 3/4 will be yellow and 1/4 green; similarly, among the 1/4 of the progeny that are wrinkled, 3/4 will be yellow and 1/4 green. The overall proportions of round yellow to round green to wrinkled yellow to wrinkled green are therefore expected to be

$$\begin{aligned} & 3/4 \times 3/4 : 3/4 \times 1/4 : 1/4 \times 3/4 : 1/4 \times 1/4 \\ & = 9/16 : 3/16 : 3/16 : 1/16 \end{aligned}$$

The observed ratio of 315 : 108 : 101 : 32 equals 9.84 : 3.38 : 3.16 : 1, which is a satisfactory fit to the 9 : 3 : 3 : 1 ratio expected from the Punnett square in Figure 3.8.

The Principle of Independent Assortment

The independent segregation of the *W*, *w* and *G*, *g* allele pairs is illustrated in Figure 3.9. What independence means is that if a gamete contains *W*, it is equally likely to contain *G* or *g*; and if a gamete contains *w*, it is equally likely to contain *G* or *g*. The implication is that the four gametes are formed in equal frequencies:

$$1/4 \text{ } WG \quad 1/4 \text{ } Wg \quad 1/4 \text{ } wG \quad 1/4 \text{ } wg$$

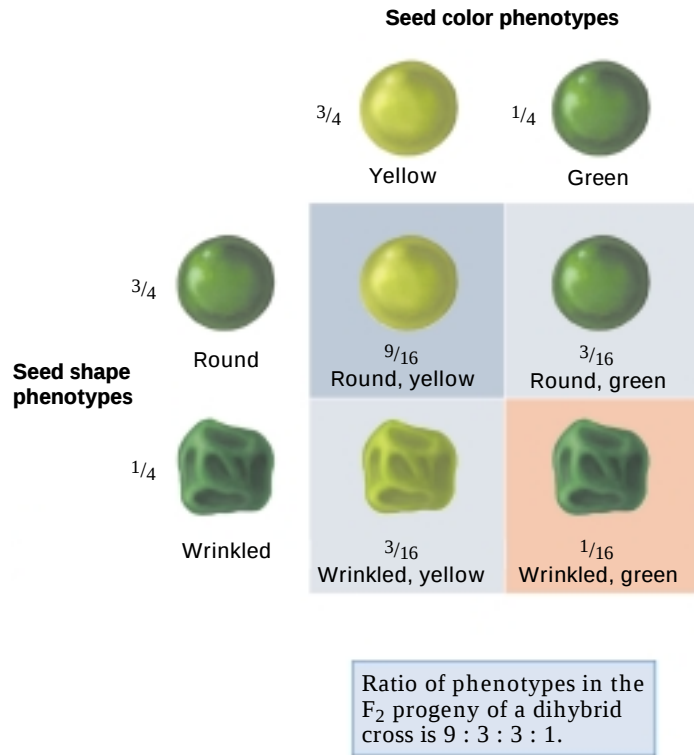


Figure 3.8 The 3 : 1 ratio of round : wrinkled, when combined at random with the 3 : 1 ratio of yellow : green, yields the 9 : 3 : 3 : 1 ratio observed in the F₂ progeny of the dihybrid cross.

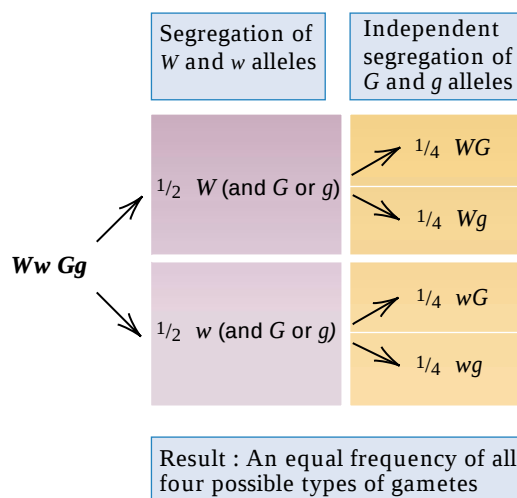


Figure 3.9 Independent segregation of the *W*, *w* and *G*, *g* allele pairs means that among each of the *W* and *w* gametic classes, the ratio of *G* : *g* is 1 : 1. Likewise, among each of the *G* and *g* gametic classes, the ratio of *W* : *w* is 1 : 1.

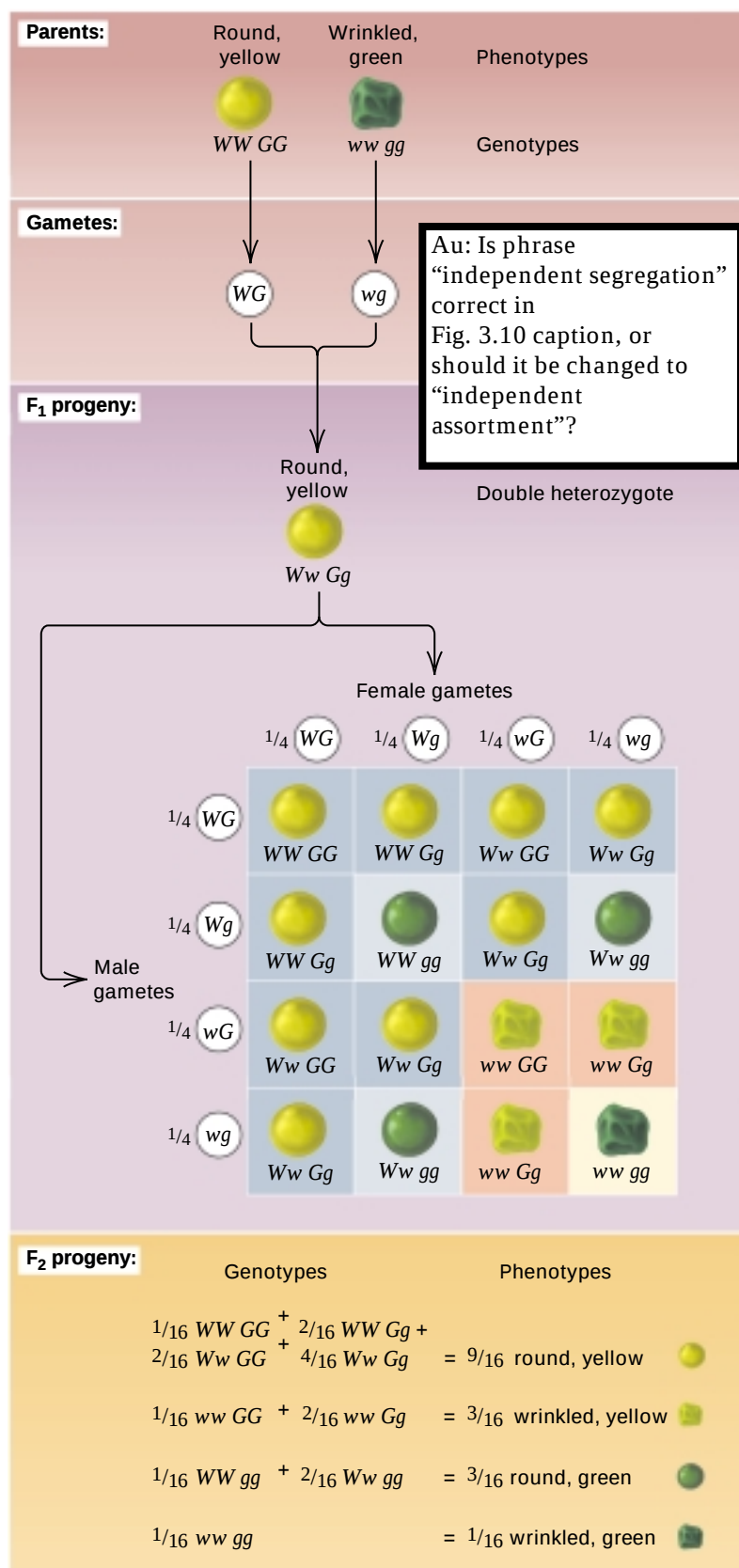


Figure 3.10 Independent segregation is the biological basis for the 9 : 3 : 3 : 1 ratio of F₂ phenotypes resulting from a dihybrid cross.

The result of independent assortment when the four types of gametes combine at random to form the zygotes of the next generation is shown in Figure 3.10. Note that the expected ratio of phenotypes among the F₂ progeny is

$$9 : 3 : 3 : 1$$

However, as the Punnett square also shows, the ratio of genotypes in the F₂ generation is more complex; it is

$$1 : 2 : 1 : 2 : 4 : 2 : 1 : 2 : 1$$

The reason for this ratio is shown in Figure 3.11. Among seeds that have the WW genotype, the ratio of

$$GG : Gg : gg \text{ equals } 1 : 2 : 1$$

Among seeds that have the Ww genotype, the ratio is

$$2 : 4 : 2$$

(which is a 1 : 2 : 1 ratio multiplied by 2 because there are twice as many Ww genotypes as either WW or ww). And among seeds that have the ww genotype, the ratio of

$$GG : Gg : gg \text{ equals } 1 : 2 : 1$$

The phenotypes of the seeds are shown beneath the genotypes, and the combined phenotypic ratio is

$$9 : 3 : 3 : 1$$

Figure 3.11 also shows that among seeds that are GG, the ratio of

$$WW : Ww : ww \text{ equals } 1 : 2 : 1$$

among seeds that are Gg, it is

$$2 : 4 : 2$$

and among seeds that are gg, it is

$$1 : 2 : 1$$

Therefore, the independent segregation means that among each of the possible genotypes formed by one pair of alleles, the ratio of homozygous dominant to heterozygous to homozygous recessive is 1 : 2 : 1 for the other independent pair of alleles.

The principle of independent segregation of two pairs of alleles in different chromosomes (or located sufficiently far apart in the same chromosome) has come to be known as the principle of independent as-

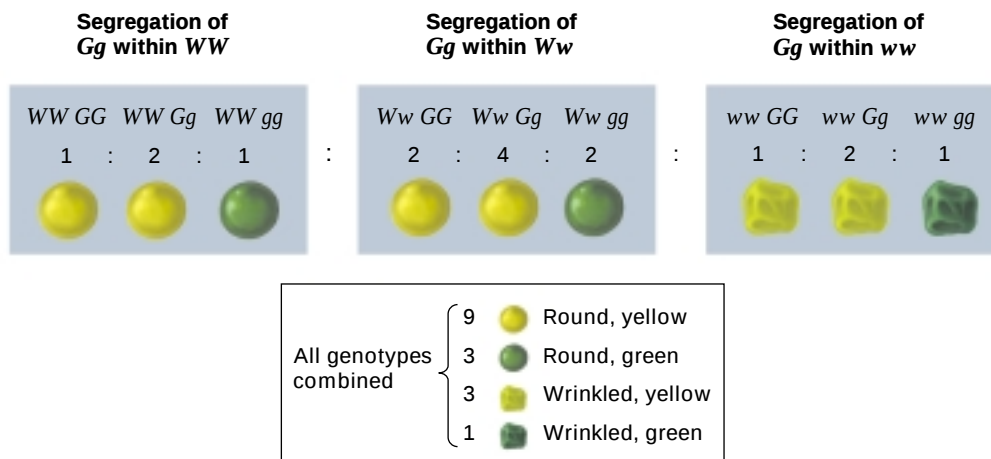


Figure 3.11 Genotypes and phenotypes of the F_2 progeny of the dihybrid cross for seed shape and seed color.

sortment. It is also sometimes referred to as *Mendel's second law*.

The Principle of Independent Assortment:

Segregation of the members of any pair of alleles is independent of the segregation of other pairs in the formation of reproductive cells.

Although the principle of independent assortment is fundamental in Mendelian genetics, the phenomenon of linkage, caused by proximity of genes in the same chromosome, is an important exception.

As with the one-gene testcross, in a two-gene testcross the ratio of progeny phenotypes is a direct demonstration of the ratio of gametes produced by the doubly heterozygous parent. In the actual cross, Mendel obtained 55 round yellow, 51 round green, 49 wrinkled yellow, and 53 wrinkled green, which is in good agreement with the predicted 1 : 1 : 1 : 1 ratio.

Au: Proof asks if it is OK that term in Key Terms list is "independent assortment." Or should term in list be changed to "principle of independent assortment"? Or should term be made bold elsewhere in text?

The Testcross with Unlinked Genes

Genes that show independent assortment are said to be **unlinked**. The hypothesis of independent assortment can be tested directly in a testcross with the double homozygous recessive:

$$Ww Gg \times ww gg$$

The result of the testcross is shown in **Figure 3.12**. Because plants with doubly heterozygous genotypes produce four types of gametes— $W G$, $W g$, $w G$, and $w g$ —in equal frequencies, whereas the plants with $ww gg$ genotypes produce only $w g$ gametes, the possible progeny genotypes are $Ww Gg$, $Ww gg$, $ww Gg$, and $ww gg$, and these are expected in equal frequencies. Because of the dominance relations— W over w and G over g —the progeny phenotypes are expected to be round yellow, round green, wrinkled yellow, and wrinkled green in a ratio of

$$1 : 1 : 1 : 1$$

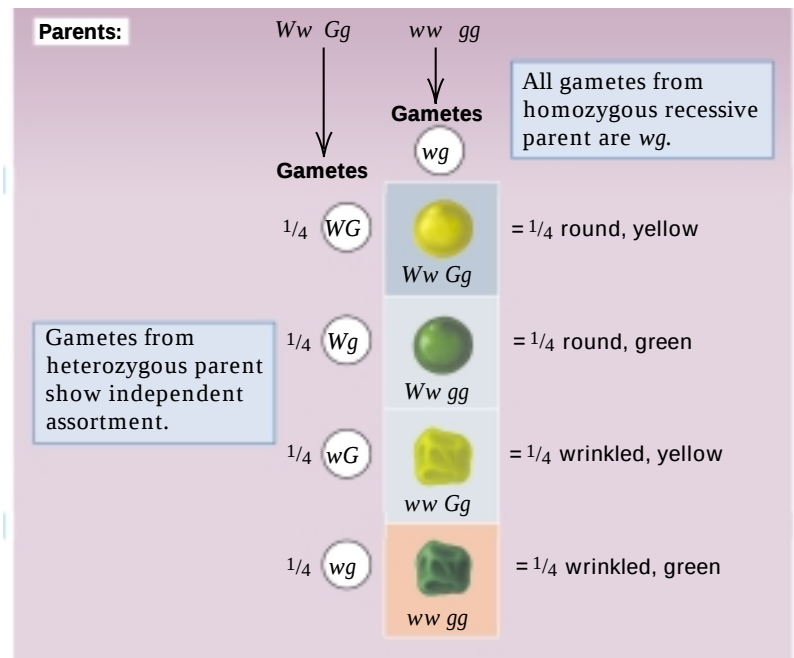


Figure 3.12 Genotypes and phenotypes resulting from a testcross of the $Ww Gg$ double heterozygote.

The Big Experiment

By analogy with the independent segregation of two genes illustrated in Figure 3.8, one might expect that independent segregation of three genes would produce F₂ progeny in which combinations of the phenotypes are given by successive terms in the multiplication of [(3/4) + (1/4)]³, which, when multiplied out, yields 27 : 9 : 9 : 9 : 3 : 3 : 3 : 1. A specific example is shown in Figure 3.13. The allele pairs are W, w for round versus wrinkled seeds; G, g for yellow versus green seeds; and P, p for purple versus white flowers. (The alleles indicated with the uppercase letters are dominant.) With independent segregation, among the F₂ progeny the most frequent phenotype (27/64) has the dominant form of all three traits, the next most frequent (9/64) has the dominant form of two of the traits, the next most frequent (3/64) has the dominant form of only one trait, and the least frequent (1/64) is the triple recessive. Observe that if you consider any one of the traits and ignore the other two, then the ratio of phenotypes is 3 : 1; and if you consider any two of the traits, then the ratio of phenotypes is 9 : 3 : 3 : 1. This means that all of the possible one- and two-gene independent segregations are present in the overall three-gene independent segregation.

Mendel tested this hypothesis, too, but by this time he was complaining of the amount of work the experiment entailed, noting that “of all the experiments, it required the most time and effort.” The result of the experiment is shown in Figure 3.14. From top to bottom, the three Punnett squares show the segregation of W and w from G and g in the genotypes PP, Pp, and pp. In each box, the number in red is the expected number of plants of each genotype, assuming independent assortment, and the number in black is the observed number of each genotype of plant. The excellent agreement confirmed what Mendel regarded as the main conclusion of all of his experiments:

“Pea hybrids form germinal and pollen cells that in their composition correspond in equal numbers to all the constant forms resulting from the combination of traits united through fertilization.”

In this admittedly somewhat turgid sentence, Mendel incorporated both segregation and independent assortment. In modern terms, what he means is that with independent segregation, the gametes produced by any hybrid plant consist of equal numbers of all possible combinations of the alleles that are heterozygous. For example,

Figure 3.13 With independent assortment, the expected ratio of phenotypes in a trihybrid cross is obtained by multiplying the three independent 3 : 1 ratios of the dominant and recessive phenotypes. A dash used in a genotype symbol indicates that either the dominant or the recessive allele is present; for example, W— refers collectively to the genotypes WW and Ww. (The expected numbers total 640 rather than 639 because of round-off error.)

(3/4 W— + 1/4 ww) × (3/4 G— + 1/4 gg) × (3/4 P— + 1/4 pp)					
				Observed number	Expected number
27/64	W— G— P—	Round, yellow, purple		269	270
9/64	W— G— pp	Round, yellow, white		98	90
9/64	W— gg P—	Round, green, purple		86	90
9/64	ww G— P—	Wrinkled, yellow, purple		88	90
3/64	W— gg pp	Round, green, white		27	30
3/64	ww G— pp	Wrinkled, yellow, white		34	30
3/64	ww gg P—	Wrinkled, green, purple		30	30
1/64	ww gg pp	Wrinkled, green, white		7	10
● For any one gene, the ratio of phenotypes is 48 : 16 = 3 : 1				● For any pair of genes, the ratio of phenotypes is 36 : 12 : 12 : 4 = 9 : 3 : 3 : 1	

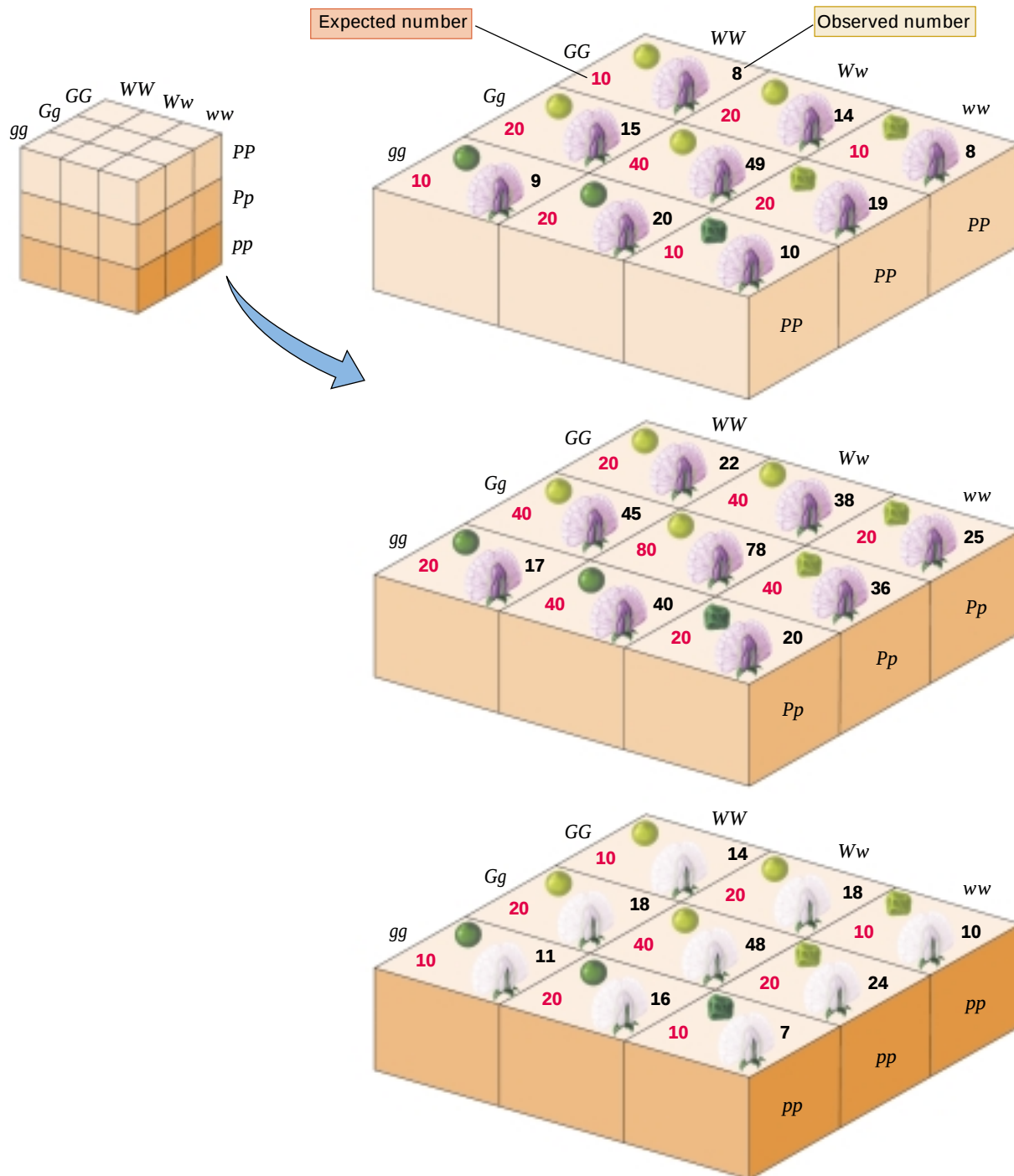


Figure 3.14 Progeny observed in the F_2 generation of a trihybrid cross with the allelic pairs W , w and G , g and P , p . In each box, the red entry is the expected number and the black entry is the

observed. Note that each gene, by itself, yields a 1 : 2 : 1 ratio of genotypes and that each pair of genes yields a 1 : 2 : 1 : 2 : 4 : 2 : 1 : 2 : 1 ratio of genotypes.

the cross $WW\ gg \times ww\ GG$ produces F_1 progeny of genotype $Ww\ Gg$, which yields the gametes WG , Wg , wG , and wg in equal numbers. Segregation is illustrated by the 1

: 1 ratio of W : w and G : g gametes, and independent assortment is illustrated by the equal numbers of WG , Wg , wG , and wg gametes.

3.4 Human Pedigree Analysis

Large deviations from expected genetic ratios are often found in individual human families and in domesticated large animals because of the relatively small number of progeny. The effects of segregation are nevertheless evident upon examination of the phenotypes among several generations of related individuals. A diagram of a family tree showing the phenotype of each individual among a group of relatives is a **pedigree**. In this section we introduce basic concepts in pedigree analysis.

Characteristics of Dominant and Recessive Inheritance

Figure 3.15 defines the standard symbols used in depicting human pedigrees. Females are represented by circles and males by squares. (A diamond is used if the sex is unknown—as, for example, in a miscarriage.) Persons with the phenotype of interest are indicated by colored or shaded symbols. For recessive alleles, heterozygous carriers are depicted with half-filled symbols. A mating between a female and a male is indicated by joining their symbols with a horizontal line, which is connected vertically to a second horizontal line, below, that connects the

symbols for their offspring. The offspring within a sibship, called **siblings** or **sibs**, are represented from left to right in order of their birth.

A pedigree for the trait *Huntington disease*, caused by a dominant mutation, is shown in Figure 3.16. The numbers in the pedigree are for convenience in referring to particular persons. The successive generations are designated by Roman numerals. Within any generation, all of the persons are numbered consecutively from left to right. The pedigree starts with the woman I-1 and the man I-2. The man has Huntington disease, which is a progressive nerve degeneration that usually begins about middle age. It results in severe physical and mental disability and then death. The dominant allele, *HD*, that causes Huntington disease is rare. All affected persons in the pedigree have the heterozygous genotype *HD hd*, whereas nonaffected persons have the homozygous normal genotype *hd hd*. The disease has complete penetrance. The **penetrance** of a genetic disorder is the proportion of individuals with the at-risk genotype who actually express the trait; *complete penetrance* means the trait is expressed in 100 percent of persons with that genotype. The pedigree demonstrates the following characteristic features of inheri-

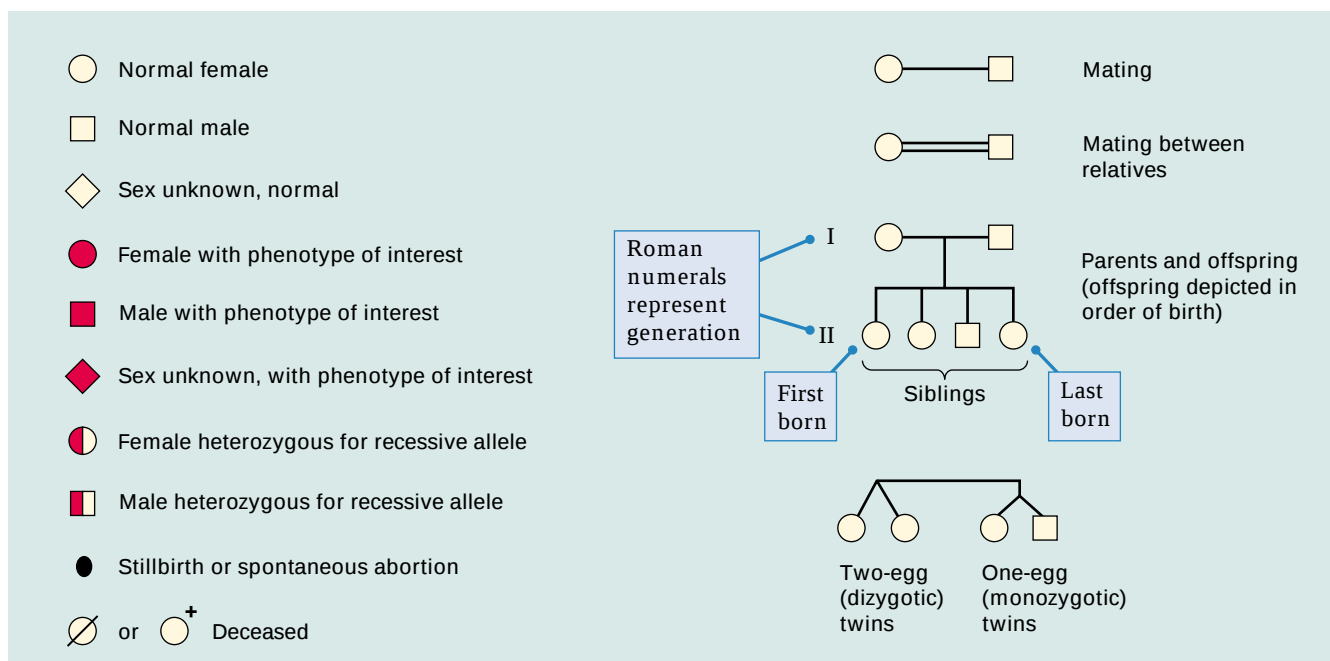


Figure 3.15 Conventional symbols used in depicting human pedigrees.

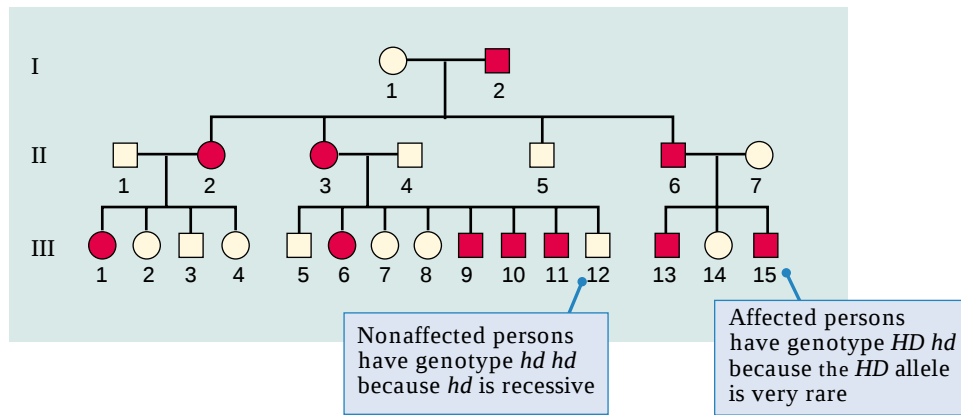


Figure 3.16 Pedigree of a human family showing the inheritance of the dominant gene for Huntington disease. Females and males are represented by circles and squares, respectively. Red symbols indicate persons affected with the disease.

tance due to a rare dominant allele with complete penetrance.

1. Females and males are equally likely to be affected.
2. Affected offspring have one affected parent (except for rare new mutations), and the affected parent is equally likely to be the mother or the father.
3. On average, half of the individuals in sibships with an affected parent are affected.

A pedigree for a trait due to a homozygous recessive allele is shown in **Figure 3.17**. The trait is *albinism*, absence of pigment in the skin, hair, and iris of the eyes. This pedigree illustrates characteristics of inheritance due to a rare recessive allele with complete penetrance:

1. Females and males are equally likely to be affected.
2. Affected individuals, if they reproduce, usually have unaffected progeny.
3. Most affected individuals have unaffected parents.
4. The parents of affected individuals are often relatives.
5. Among siblings of affected individuals, the proportion affected is approximately 25 percent.

With rare recessive inheritance, the mates of homozygous affected persons are usually homozygous for the normal allele, so all of the offspring will be heterozygous and not affected. Heterozygous **carriers** of the mutant allele are considerably more common than homozygous affected individuals, because it is more likely that a person will inherit only one copy of a rare mutant allele than two copies. Most homozygous recessive genotypes therefore result from

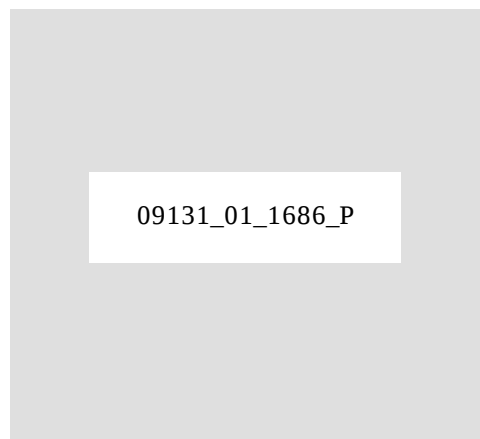


Figure 3.17 Pedigree of albinism. With recessive inheritance, affected persons (filled symbols) often have unaffected parents. The double horizontal line indicates a mating between relatives—in this case, first cousins.

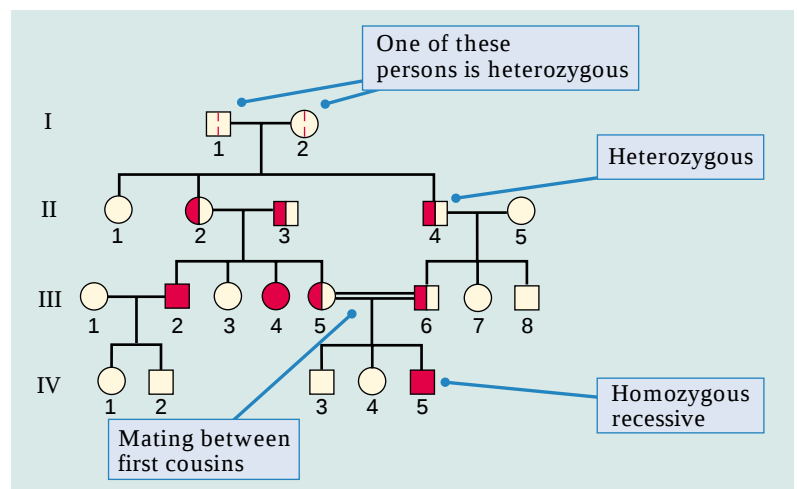


Figure 3.17 Pedigree of albinism. With recessive inheritance, affected persons (filled symbols) often have unaffected parents. The double horizontal line indicates a mating between relatives—in this case, first cousins.

matings between carriers (heterozygous $\times$ heterozygous), in which each offspring has a 1/4 chance of being affected. Another important feature of rare recessive inheritance is that the parents of affected individuals are often related (**consanguineous**). A mating between relatives is indicated with a double line connecting the partners, as for the first-cousin mating in Figure 3.17. Mating between relatives is important for recessive alleles to become homozygous, because when a recessive allele is rare, it is more likely to become homozygous through inheritance from a common ancestor than from parents who are completely unrelated. The reason is that the carrier of a rare allele may have many descendants who are also carriers. If two of these carriers should mate with each other (for example, in a first-cousin mating), then the hidden recessive allele can become homozygous with a probability of 1/4. Mating between relatives constitutes *inbreeding*, and the consequences of inbreeding are discussed further in Chapter 17. Because an affected individual indicates that the parents are heterozygous carriers, the expected proportion of affected individuals among the siblings is approximately 25 percent, but the exact value depends on the details of how affected individuals are identified and included in the database.

- Most genes that cause genetic diseases are rare, so they are observed in only a small number of families.
- Many genes of interest in human genetics are recessive, so they are not detected in heterozygous genotypes.
- The number of offspring per human family is relatively small, so segregation cannot usually be detected in single sibships.
- The human geneticist cannot perform testcrosses or backcrosses, because human matings are not manipulated by an experimenter.

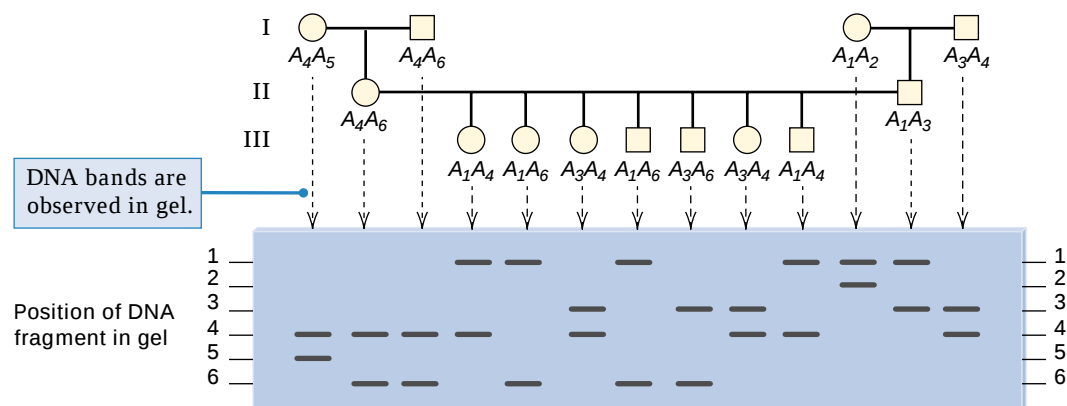
Because techniques for manipulating DNA allow direct access to the DNA, modern genetic studies of human pedigrees are carried out primarily using genetic markers present in the DNA itself, rather than through the phenotypes produced by mutant genes. The human genome includes one single-nucleotide polymorphism per 500–3000 bp, depending on the region being studied. This means that two randomly chosen human genomes differ at 1–6 million nucleotide positions.

Various types of DNA polymorphisms were discussed in Chapter 2, along with the methods by which they are detected and studied. An example of a DNA polymorphism segregating in a three-generation human pedigree is shown in Figure 3.18. The type of polymorphism is a *simple tandem repeat polymorphism (STRP)*, in which each allele differs in size according to the number of copies it contains of a short DNA sequence repeated in tandem. The differences in size are detected by electrophoresis after amplification of the region by PCR. STRP markers usually have as many as 20 codominant alleles, and the majority of in-

Molecular Markers in Human Pedigrees

Before the advent of molecular methods, there were many practical obstacles to the study of human genetics.

Figure 3.18 Human pedigree showing segregation of VNTR alleles. Six alleles (A_1 – A_6) are present in the pedigree, but any one person can have only one allele (if homozygous) or two alleles (if heterozygous).



dividuals are heterozygous for two different alleles. In the example in Figure 3.18, each of the parents is heterozygous, as are all of the children. More than 5000 genetic markers of this type have been identified in the human genome, each heterozygous in an average of 70 percent of individual genotypes.

Six alleles are depicted in Figure 3.18, denoted by A_1 through A_6 . In the gel, the numbers of the bands correspond to the subscripts of the alleles. The mating in generation II is between two heterozygous genotypes: $A_4A_6 \times A_1A_3$. Because of segregation in each parent, four genotypes are possible among the offspring (A_4A_1 , A_4A_3 , A_6A_1 , and A_6A_3); these would conventionally be written with the smaller subscript first, as A_1A_4 , A_3A_4 , A_1A_6 , and A_3A_6 . With random fertilization the offspring genotypes are equally likely, as may be verified from a Punnett square for the mating. Figure 3.18 illustrates some of the principal advantages of multiple, codominant alleles for human pedigree analysis: (1) Heterozygous genotypes can be distinguished from homozygous genotypes. (2) Many individuals in the population are heterozygous, and so many matings are informative in regard to segregation. (3) Each segregating genetic marker yields up to four distinguishable offspring genotypes.

3.5 Pedigrees and Probability

Genetic analysis benefits from large numbers of progeny because these numbers determine the degree to which observed ratios of genotypes and phenotypes fit their expected values. Greater statistical variation occurs in smaller numbers—and therefore potentially greater deviations from the expected values. For example, among the seeds of individual pea plants in which a 3 : 1 ratio was expected, Mendel observed ratios ranging from 1.85 : 1 to 4.85 : 1. The large variation results from the relatively small number of seeds per plant, which averaged about 34 in these experiments. Taken together, the seeds yielded a phenotypic ratio of 3.08 : 1, which not only fits 3 : 1 but is actually a better fit than that observed in any of the individual plants. Because of statistical variation in small

numbers, a working knowledge of probability is basic to understanding genetic transmission. In the first place, each event of fertilization represents a chance combination of alleles present in the parental gametes. In the second place, the proportions of the different types of offspring obtained from a cross are the cumulative result of numerous independent events of fertilization.

In the analysis of genetic crosses, the probability of a particular outcome of a fertilization event may be considered as equivalent to the proportion of times that this outcome is expected to be realized in numerous repeated trials. The reverse is also true: The proportion of times that an outcome is expected to be realized in numerous repeated trials is equivalent to the probability that it is realized in a single trial. To take a specific example using the pedigree in Figure 3.18, the mating at the upper left is $A_4A_5 \times A_4A_6$. Among a large number of offspring from such a mating, the expected proportion of A_4A_6 genotypes is 1/4. Equivalently, we could say that any offspring whose genotype is unknown has a probability of genotype A_4A_6 equal to 1/4. However, the genotype of the female offspring shown is already given as A_4A_6 , so relative to this individual, the situation has become a certainty: The probability that her genotype is A_4A_6 equals 1, and the probability that her genotype is anything else equals 0.

To evaluate the probability of a genetic event usually requires an understanding of the mechanism of inheritance and knowledge of the particular cross. For example, in asserting that the mating I-1 $\times$ I-2 in Figure 3.18 has a probability of yielding an A_4A_6 offspring equal to 1/4, we needed to take the parental genotypes into account. If the mating were $A_4A_6 \times A_4A_6$, then the probability of an A_4A_6 offspring would be 1/2 rather than 1/4; if the mating were $A_4A_4 \times A_6A_6$, then the probability of an A_4A_6 offspring would be 1.

In many genetic crosses, the possible outcomes of fertilization are equally likely. Suppose that there are n possible outcomes, each as likely as any other, and that in m of these, a particular outcome of interest is realized; then the probability of the outcome of interest is m/n . In the language of probability, a possible outcome of interest is typically called an *event*. As an example,

Huntington disease (HD) is a progressive neurodegenerative disorder characterized by motor disturbance, cognitive loss, and psychiatric manifestations. It is inherited in an autosomal dominant fashion and affects approximately 1 in 10,000 individuals in most populations of European origin. The hallmark of HD is a distinctive choreic [jerky] movement disorder that typically has a subtle, insidious onset in the fourth to fifth decade of life and gradually worsens over a course of 10 to 20 years until death. . . . The

genetic defect causing HD was assigned to chromosome 4 in one of the first successful linkage analyses using DNA markers in humans. Since that time, we have pursued an approach to isolating and characterizing the HD gene based on progressively refining its localization. . . . [We have found that a] 500-kb segment is the most likely site of the genetic defect. [The abbreviation kb stands for kilobase pairs; 1 kb equals 1000 base pairs.] Within this region, we have identified a large gene, spanning approximately 210 kb, that encodes a previously undescribed protein. The reading frame contains a polymorphic (CAG)_n trinucleotide repeat with at least 17 alleles in the normal population, varying from 11 to 34 CAG copies. On HD chromo-

somes, the length of the trinucleotide repeat is substantially increased. . . . Elon-

We consider it of the utmost importance that the current internationally accepted guidelines and counseling protocols for testing people at risk continue to be observed, and that samples from unaffected relatives should not be tested inadvertently or without full consent.

gation of a trinucleotide repeat sequence has been implicated previously as the cause of three quite different human disorders, the fragile-X syndrome, myotonic dystrophy, and spino-bulbar muscular atrophy. . . . It can be expected that the capacity to monitor directly the size of the trinucleotide repeat in individuals “at risk” for HD will revolutionize testing for the disorder. . . . We con-

sider it of the utmost importance that the current internationally accepted guidelines and counseling protocols for testing people at risk continue to be observed, and that samples from unaffected relatives should not be tested inadvertently or without full consent. . . . With the mystery of the genetic basis of HD apparently solved, [it opens] the next challenges in the effort to understand and to treat this devastating disorder.

Source: *Cell* 72: 971–983

of interest, and different types of groupings may have different probabilities associated with them.

Mutually Exclusive Possibilities

Sometimes an outcome of interest includes two or more different possibilities. The offspring of the mating $A_4A_5 \times A_4A_6$ again provides a convenient example. As we have seen, if the offspring are classified as either homozygous or heterozygous, then the outcome “homozygous” consists of just one possibility (A_4A_4) and the outcome “heterozygous” consists of three possibilities (A_4A_6 , A_4A_5 , and A_5A_6). Furthermore, if an individual genotype is any one of the three heterozygous genotypes, it cannot at the same time be any of the other heterozygous genotypes. In other words, only one genotype can be realized in any one organism, so the realization of one possibility precludes the realization of another in the same organism. Events that exclude each other in this manner are said to be *mutually exclusive*. When events are mutually exclu-

sive, their probabilities are combined according to the addition rule.

Addition Rule: The probability of the realization of one or the other of two mutually exclusive events, A or B, is the sum of their separate probabilities.

In symbols, if we use *Prob* to mean *probability*, then the addition rule for two possible outcomes is written

$$\text{Prob}\{A \text{ or } B\} = \text{Prob}\{A\} + \text{Prob}\{B\}$$

Application of the addition rule is straightforward in the above examples of homozygous versus heterozygous genotypes, but there are three possible outcomes instead of two. Because the offspring genotypes are equally likely, the probability of a heterozygous offspring equals $\text{Prob}\{A_4A_6, A_4A_5, \text{ or } A_5A_6\} = \text{Prob}\{A_4A_6\} + \text{Prob}\{A_4A_5\} + \text{Prob}\{A_5A_6\} = 1/4 + 1/4 + 1/4 = 3/4$. Because 3/4 is the probability of an individual offspring being heterozygous, it is also the expected proportion of heterozygous genotypes among a large number of progeny.

Independent Possibilities

Events that are not mutually exclusive may be *independent*, which means that the realization of one outcome has no influence on the possible realization of any others. We have already seen an example in the independent segregation of round versus wrinkled as against yellow versus green peas (Figure 3.9). The principle is that when the possible outcomes of an experiment or observation are independent, the probability that they are realized together is obtained by multiplication.

Successive offspring from a cross are also independent events, which means that the genotypes of early progeny have no influence on the probabilities in later progeny (Figure 3.19). The independence of successive offspring contradicts the widespread belief that in each human family, the ratio of girls to boys must “even out” at approximately 1 : 1. According to this reasoning, a family with four girls would be more likely to have

a boy the next time around. But this belief is supported neither by theory nor by actual data on the sex ratio in human sibships. The data indicate that human families are equally likely to have a girl or a boy on any birth, irrespective of the sex distribution in previous births. Although statistics guarantees that the sex ratio will balance out when averaged over a very large number of sibships, this does not imply that it will equalize in any individual sibship. To be concrete, among families in which there are five children, sibships consisting of five boys balance those consisting of five girls, yielding an overall sex ratio of 1 : 1; nevertheless, both of these sibships have an unusual sex ratio.

When events are independent (such as independent traits or successive offspring from a cross), the probabilities are combined by means of the multiplication rule.

Multiplication Rule: The probability of two independent events, A and B, being realized simultaneously is given by the product of their separate probabilities.

In symbols, the multiplication rule is written

$$\text{Prob}\{A \text{ and } B\} = \text{Prob}\{A\} \cdot \text{Prob}\{B\}$$

The multiplication rule can be used to answer questions such as this: For two offspring from the mating $A_4A_5 \times A_4A_6$, what is the probability of one homozygous genotype and one heterozygous genotype? The birth order homozygous–heterozygous has probability $1/4 \times 3/4$, and the birth order heterozygous–homozygous has probability $3/4 \times 1/4$. These possibilities are mutually exclusive, so the overall probability is

$$1/4 \times 3/4 + 3/4 \times 1/4 = 2(1/4)(3/4) = 3/8$$

Note that the addition rule was used twice in solving this problem, once to calculate the probability of a heterozygous offspring and again to calculate the probability of both birth orders. The multiplication rule can also be used to calculate the probability of a specific genotype among the progeny of a complex cross. For example, if a quadruple heterozygous genotype $Aa Bb Cc Dd$ is self-fertilized, what is the probability of a quadruple heterozygous offspring, $Aa Bb Cc Dd$? Assuming independent assortment, the answer is $(1/2)(1/2)(1/2)(1/2) = (1/2)^4 = 1/16$, because each of the inde-

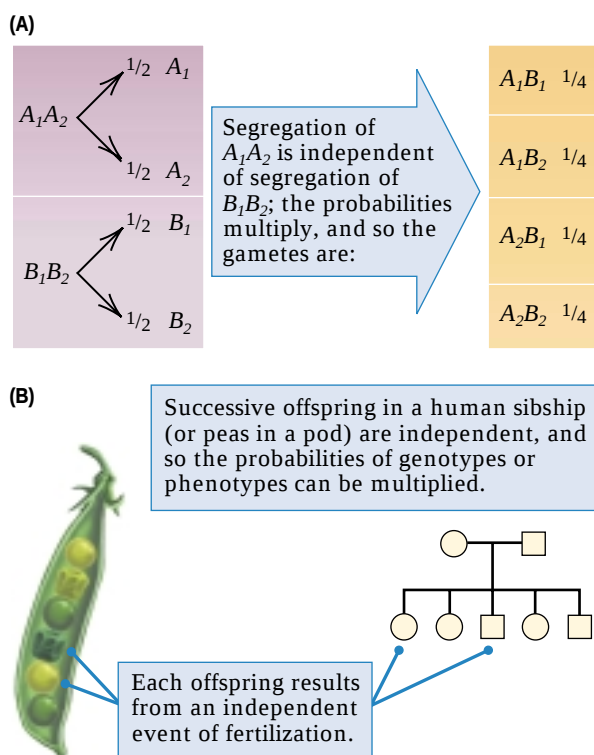


Figure 3.19 In genetics, two important types of independence are (A) independent segregation of alleles that show independent assortment and (B) independent fertilizations resulting in successive offspring. In these cases, the probabilities of the individual outcomes of segregation or fertilization are multiplied to obtain the overall probability.

pendent genes yields a heterozygous offspring with probability $1/2$.

3.6 Incomplete Dominance and Epistasis

Dominance and codominance are not the only possibilities for pairs of alleles. There are situations of **incomplete dominance**, in which the phenotype of the heterozygous genotype is intermediate between the phenotypes of the homozygous genotypes. A classic example of incomplete dominance concerns flower color in the snapdragon *Antirrhinum majus* (Figure 3.20). In wildtype flowers, a red type of anthocyanin pigment is formed by a sequence of enzymatic reactions. A wildtype enzyme, encoded by the *I* allele, is limiting to the rate of the overall reaction, so the amount of red pigment is determined by the amount of enzyme that the *I* allele produces. The alternative *i* allele codes for an inactive enzyme, and *ii* flowers are ivory in color. Because the amount of the critical enzyme is reduced in *Ii* heterozygotes, the amount of red pigment in the flowers is reduced also, and the effect of the dilution is to make the flowers pink. A cross between plants differing in flower color therefore gives direct phenotypic evidence of segregation (Figure 3.20). The cross *II* (red) $\times$ *ii* (ivory) yields F_1 plants with genotype *Ii* and pink flowers. In the F_2 progeny obtained by self-pollination of the F_1 hybrids, one experiment resulted in 22 plants with red flowers, 52 with pink flowers, and 23 with ivory flowers, which fits the expected ratio of 1 : 2 : 1.

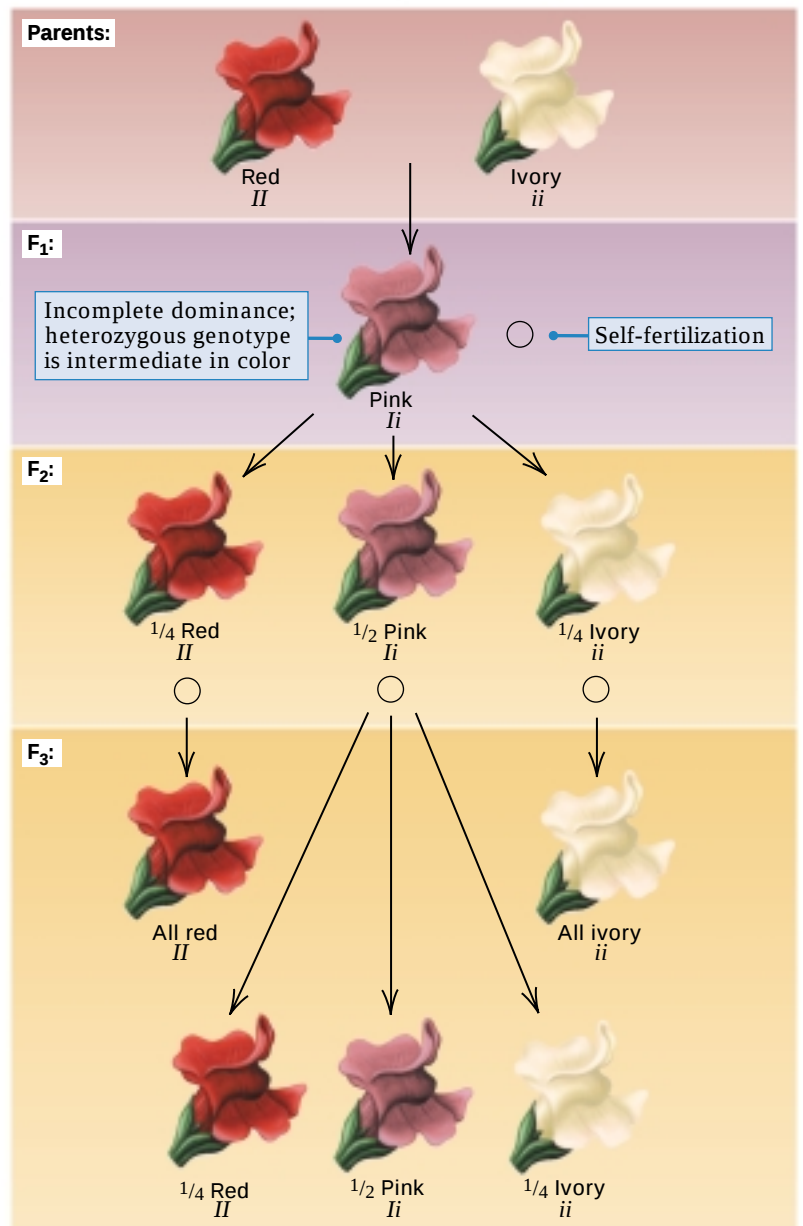


Figure 3.20 Red versus white flower color in snapdragons shows no dominance.

Multiple Alleles

The occurrence of **multiple alleles** is exemplified by the alleles A_1 – A_6 of the STRP marker in the human pedigree in Figure 3.18. Multiple alleles are relatively common in natural populations and, as in this example, can be detected most easily by molecular methods. In the DNA of a gene, each nucleotide can be A, T, G, or C, so a gene of n nucleotides can theoretically mutate at any of the positions to any of the three other nucleotides. The number of possible single-nucleotide differences in a

gene of length n is therefore $3 \times n$. If $n = 5000$, for example, there are potentially 15,000 alleles (not counting any of the possibilities with more than one nucleotide substitution). Most of the potential alleles do not actually exist at any one time. Some are absent because they did not occur, others did occur but were eliminated by chance or because they were harmful, and still others are present but at such a low frequency that they remain undetected. Nevertheless, at the level of DNA sequence, most genes in

most natural populations have multiple alleles, all of which can be considered “wild-type.” Multiple wildtype alleles are useful in such applications as DNA typing because two unrelated people are unlikely to have the same genotype, especially if several different loci, each with multiple alleles, are examined. Many harmful mutations also exist in multiple forms. Recall from Chapter 1 that more than 400 mutant forms of the phenylalanine hydroxylase gene have been identified in patients with phenylketonuria. The alleles A_1 – A_6 in Figure 3.18 also illustrate that although a *population* of organisms may contain any number of alleles, any particular organism or cell may carry no more than two, and any gamete may carry no more than one.

In some cases, the multiple alleles of a gene exist merely by chance and reflect the history of mutations that have taken place in the population and the dissemination of these mutations among population subgroups by migration and interbreeding. In other cases, there are biological mechanisms that favor the maintenance of a large number of alleles. For example, genes that control self-sterility in certain flowering plants can have large numbers of allelic types. This type of self-sterility is found in species of red clover that grow wild in many pastures. The self-sterility genes prevent self-fertilization because a pollen grain can undergo pollen tube growth and fertilization only if it contains a self-sterility allele different from either of the alleles present in the flower on which it lands. In other words, a pollen grain containing an allele already present in a flower will not function on that flower. Because all pollen grains produced by a plant must contain one of the self-sterility alleles present in the plant, pollen cannot function on the same plant that produced it, and self-fertilization cannot take place. Under these conditions, any plant with a new allele has a selective advantage, because pollen that contains the new allele can fertilize all flowers except those on the same plant. Through evolution, populations of red clover have accumulated hundreds of alleles of the self-sterility gene, many of which have been isolated and their DNA sequences determined. Many of the alleles differ at multiple nucleotide sites, which implies that the alleles in the population are very old.

Human ABO Blood Groups

In a multiple allelic series, there may be different dominance relationships between different pairs of alleles. An example is found in the human ABO blood groups, which are determined by three alleles denoted I^A , I^B , and I^O . (Actually, there are two slightly different variants of the I^A allele.) The blood group of any person may be A, B, AB, or O, depending on the type of polysaccharide (polymer of sugars) present on the surface of red blood cells. One of two different polysaccharides, A or B, can be formed from a precursor molecule that is modified by the enzyme product of the I^A or the I^B allele. The gene product is a glycosyl transferase enzyme that attaches a sugar unit to the precursor (Figure 3.21). The I^A or the I^B alleles encode different forms of the enzyme with replacements at four amino acid sites; these alter the substrate specificity so that each enzyme attaches a different sugar. People of genotype $I^A I^A$ produce red blood cells having only the A polysaccharide and are said to have blood type A. Those of genotype $I^B I^B$ have red blood cells with only the B polysaccharide and have blood type B. Heterozygous $I^A I^B$ people have red cells with both the A and the B polysaccharides and have blood type AB. The third allele, I^O , encodes an enzymatically inactive protein that leaves the precursor unchanged; neither the A nor the B type of polysaccharide is produced. Homozygous $I^O I^O$ persons therefore lack both the A and the B polysaccharides and are said to have blood type O.

In this multiple allelic series, the I^A and I^B alleles are codominant: The heterozygous genotype has the characteristics of both homozygous genotypes—the presence of both the A and the B carbohydrate on the red blood cells. On the other hand, the I^O allele is recessive to both I^A and I^B . Hence, heterozygous $I^A I^O$ genotypes produce the A polysaccharide and have blood type A, and heterozygous $I^B I^O$ genotypes produce the B polysaccharide and have blood type B. The genotypes and phenotypes of the ABO blood group system are summarized in the first three columns of Table 3.3.

ABO blood groups are important in medicine because of the need for blood transfusions. A crucial feature of the ABO system is that most human blood contains

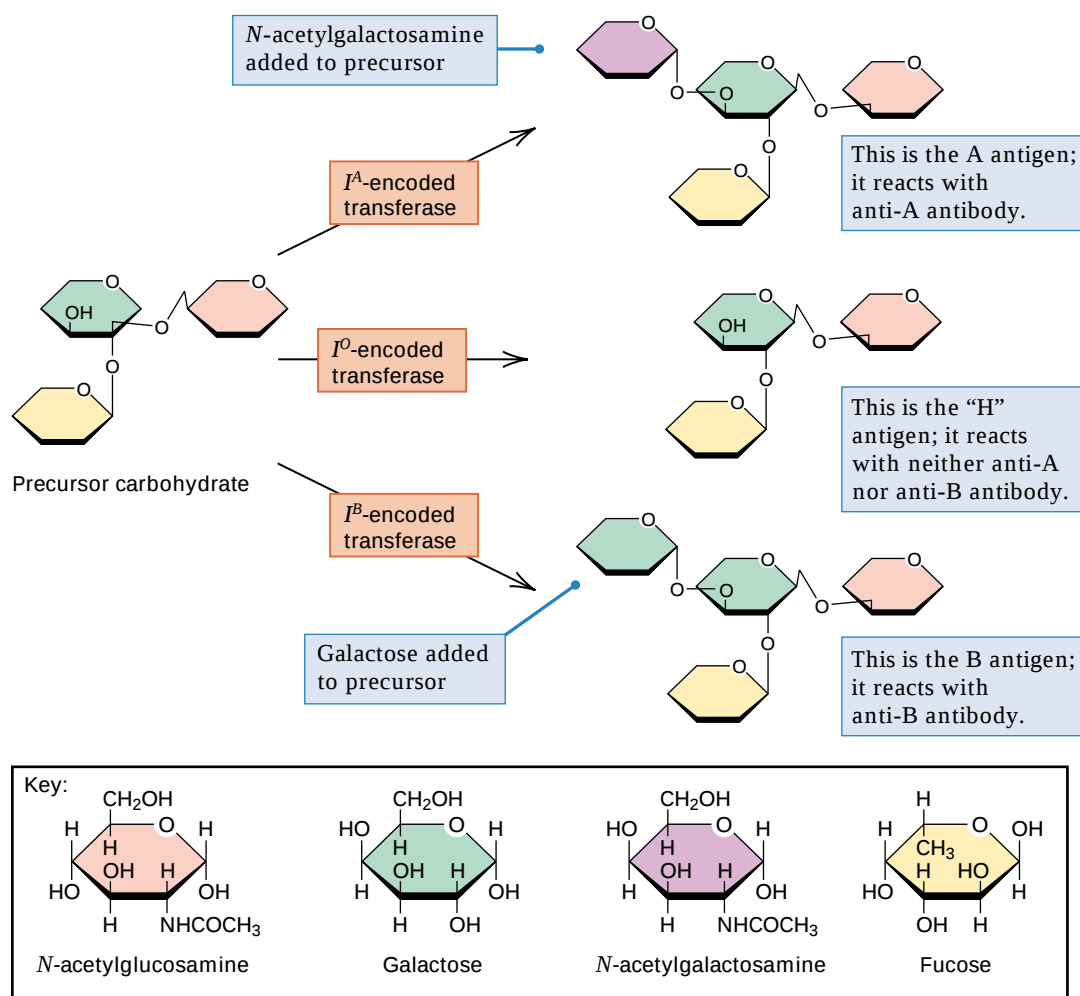


Figure 3.21 The ABO antigens on the surface of human red blood cells are carbohydrates. They are formed from a precursor carbohydrate by the action of transferase enzymes encoded by alleles of the *I* gene. Allele *I*^O codes for an inactive enzyme and leaves the precursor (called the H substance) unmodified. The *I*^A allele encodes an enzyme that adds *N*-acetylgalactosamine (purple) to the precursor. The *I*^B allele encodes an enzyme that adds galactose (green) to the precursor. The other colored sugar units are *N*-acetylglucosamine (orange) and fucose (yellow).

antibodies to either the A or the B polysaccharide. An **antibody** is a protein made by the immune system in response to a stimulating molecule called an **antigen** and is capable of binding to the antigen. An antibody is usually specific in that it recog-

nizes only one antigen. Some antibodies combine with antigen and form large molecular aggregates that may precipitate.

Antibodies act to defend against invading viruses and bacteria. Although antibodies do not normally form without prior

Table 3.3 Genetic control of the human ABO blood groups

Genotype	Antigens present on red blood cells	ABO blood group phenotype	Antibodies present in blood fluid	Blood types that can be tolerated in transfusion	Blood types that can accept blood for transfusion
$I^A I^A$	A	Type A	Anti-B	A & O	A & AB
$I^A I^O$	A	Type A	Anti-B	A & O	A & AB
$I^B I^B$	B	Type B	Anti-A	B & O	B & AB
$I^B I^O$	B	Type B	Anti-A	B & O	B & AB
$I^A I^B$	A & B	Type AB	Neither anti-A nor anti-B	A, B, AB & O	AB only
$I^O I^O$	Neither A nor B	Type O	Anti-A & anti-B	O only	A, B, AB & O

stimulation by the antigen, people capable of producing anti-A and anti-B antibodies do produce them. Production of these antibodies may be stimulated by antigens that are similar to polysaccharides A and B and that are present on the surfaces of many common bacteria. However, a mechanism called *tolerance* prevents an organism from producing antibodies against its own antigens. This mechanism ensures that A antigen or B antigen elicits antibody production only in people whose own red blood cells do not contain A or B, respectively. The end result:

People of blood type O make both anti-A and anti-B antibodies, those of blood type A make anti-B antibodies, those of blood type B make anti-A antibodies, and those of blood type AB make neither type of antibody.

The antibodies found in the blood fluid of people with each of the ABO blood types are shown in the fourth column in Table 3.3. The clinical significance of the ABO blood groups is that transfusion of blood containing A or B red-cell antigens into persons who make antibodies against them results in an agglutination reaction in which the donor red blood cells are clumped. In this reaction, the anti-A antibody agglutinates red blood cells of either blood type A or blood type AB, because both carry the A antigen (Figure 3.22). Similarly, anti-B antibody agglutinates red blood cells of either blood type B or blood type AB. When the blood cells agglutinate, many blood vessels are blocked, and the recipient of the transfusion goes into shock and may die. Incompatibility in the other

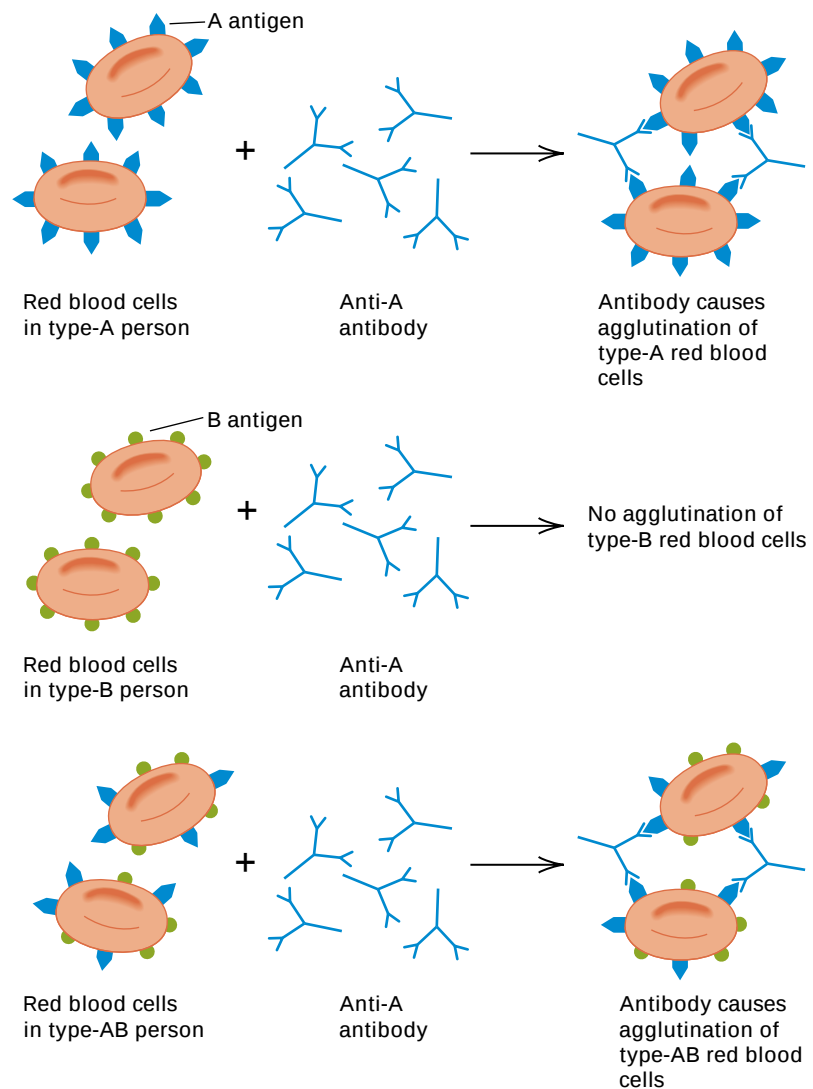


Figure 3.22 Antibody against type-A antigen agglutinates red blood cells that carry the type-A antigen, whether or not they also carry the type-B antigen. Blood fluid containing anti-A antibody agglutinates red blood cells of type A and type AB, but not red blood cells of type B or type O.

direction, in which the donor blood contains antibodies against the recipient's red blood cells, is usually acceptable because the donor's antibodies are diluted so rapidly that clumping is avoided. The types of compatible blood transfusions are shown in the last two columns of Table 3.3. Note that a person of blood type AB can receive blood from a person of any other ABO type; type AB is called a *universal recipient*. Conversely, a person of blood type O can donate blood to a person of any ABO type; type O is called a *universal donor*.

Epistasis

In Chapter 1 we saw how the products of several genes may be necessary to carry out all the steps in a biochemical pathway. In genetic crosses in which two mutations that affect different steps in a single pathway are both segregating, the typical F_2 ratio of 9 : 3 : 3 : 1 is not observed. Gene interaction that perturbs the normal Mendelian ratios is known as **epistasis**. One type of epistasis is illustrated by the interaction of the C, c and P, p allele pairs affecting flower coloration in peas. These genes encode enzymes in the biochemical pathway for the synthesis of anthocyanin pigment, and the production of anthocyanin requires the presence of at least one wildtype dominant allele of each gene. The proper way to represent this situation genetically is to write the required genotype as

$$C- P-$$

where each dash is a “blank” that may be filled with either allele of the gene. Hence $C-$ comprises the genotypes CC and Cc , and likewise $P-$ comprises the genotypes PP and Pp . All four genotypes included in the symbol $C- P-$, and only these genotypes, have purple flowers.

Figure 3.23 shows a cross between the homozygous recessive genotypes pp and cc . The phenotype of the flowers in the F_1 generation is purple because the genotype is $Cc Pp$. Self-fertilization of the F_1 plants (indicated by the encircled cross sign) results in the F_2 progeny genotypes shown in the Punnett square. Because only the $C- P-$ progeny have purple flowers, the ratio of purple flowers to white flowers in the F_2 generation is 9 : 7. The epistasis does not

change the result of independent segregation, it merely conceals the fact that the underlying ratio of the genotypes $C- P- : C- pp : cc P- : cc pp$ is 9 : 3 : 3 : 1.

For a trait determined by the interaction of two genes, each with a dominant allele,

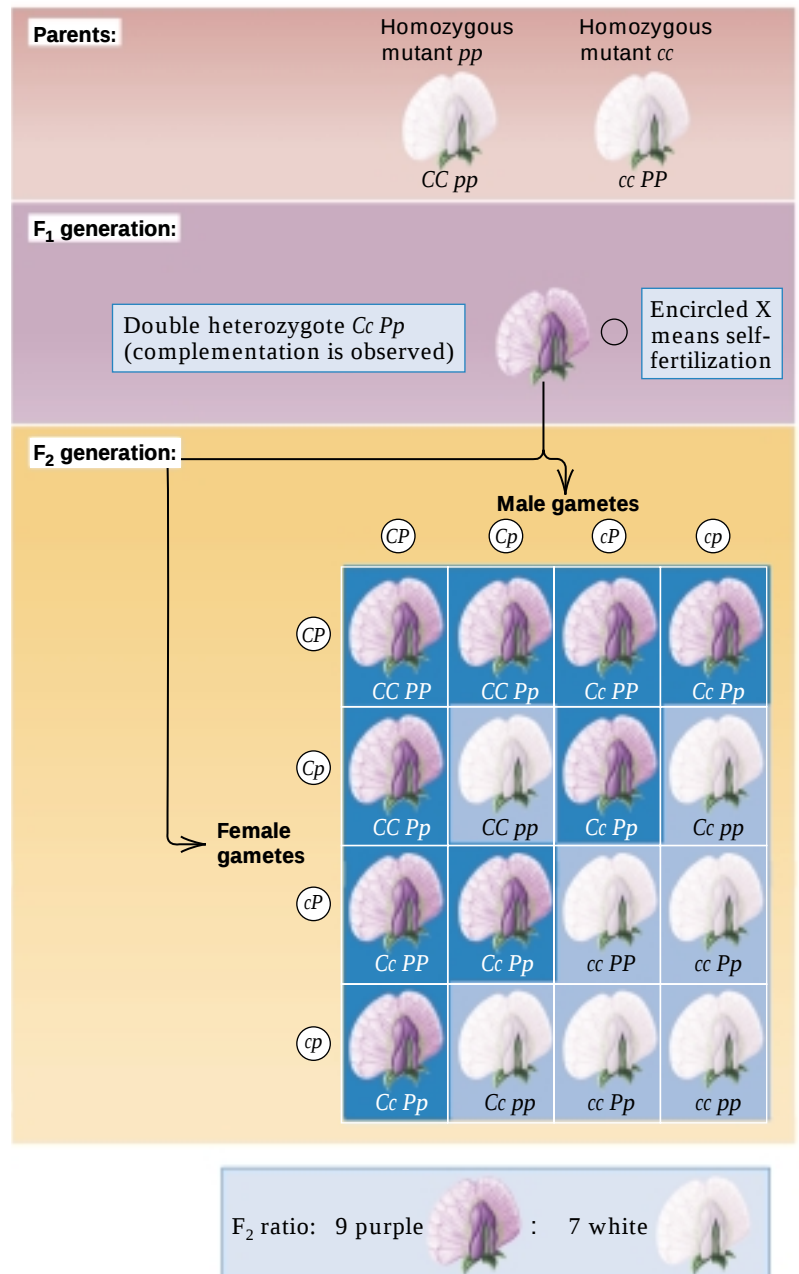


Figure 3.23 Epistasis in the determination of flower color in peas. Formation of the purple pigment requires the dominant allele of both the C and P genes. With this type of epistasis, the dihybrid F_2 ratio is modified to 9 purple : 7 white.

there are only a limited number of ways in which the 9 : 3 : 3 : 1 dihybrid ratio can be modified. The possibilities are illustrated in **Figure 3.24**. In part A are the genotypes produced in the F₂ generation by independent segregation. In the absence of epistasis, the F₂ ratio of phenotypes is 9 : 3 : 3 : 1. The possible modified ratios are shown in part B. In each row, the color coding indicates phenotypes that are indistinguishable because of epistasis, and the resulting modified ratio is given. For example, in the modified ratio at the bottom, the phenotypes of the “3 : 3 : 1” classes are indistinguishable, resulting in a 9 : 7 ratio. This is the ratio observed in the segregation of the *C, c* and *P, p* alleles in Figure 3.23, with its 9 : 7 ratio of purple to white flowers. Taking all the possible modified ratios in Figure 3.24B

together, there are nine possible dihybrid ratios when both genes show complete dominance. Examples of each of the modified ratios are known. The most frequently encountered modified ratios are 9 : 7, 12 : 3 : 1, 13 : 3, 9 : 4 : 3, and 9 : 6 : 1. The types of epistasis that result in these modified ratios are illustrated in the following examples, which are taken from a variety of organisms. Other examples can be found in the problems at the end of the chapter.

9 : 7 is observed when a homozygous recessive mutation in either or both of two different genes results in the same mutant phenotype, as in Figure 3.23.

12 : 3 : 1 results when the presence of a dominant allele at one locus masks the

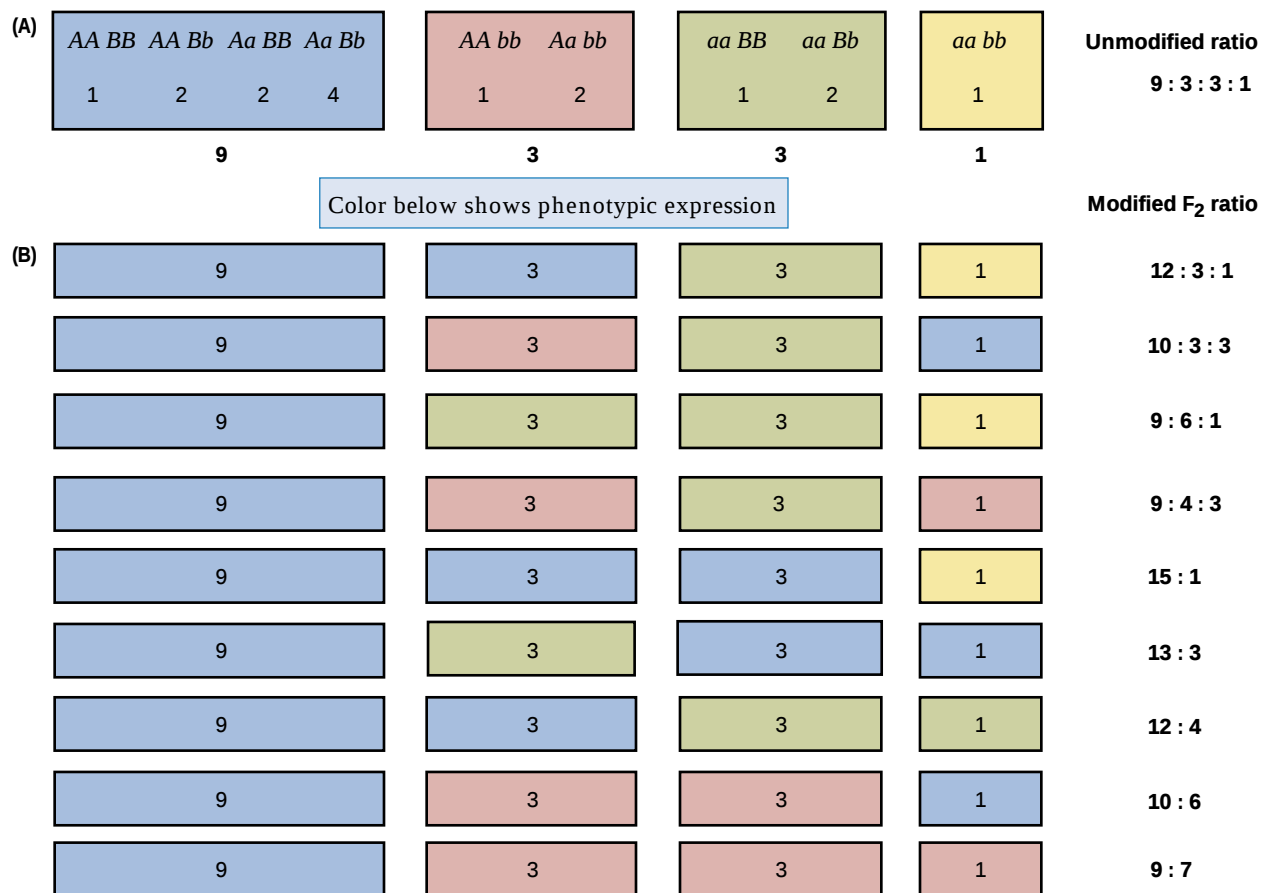


Figure 3.24 Modified F₂ dihybrid ratios. (A) The F₂ genotypes of two independently assorting genes with complete dominance result in a 9 : 3 : 3 : 1 ratio of phenotypes, provided there is no interaction (epistasis) between the genes. (B) If there is epistasis that renders two or more of the phenotypes indistinguishable (indicated by the colors), then the F₂ ratio is modified.

genotype at a different locus, such as the $A-$ genotype rendering the $B-$ and bb genotypes indistinguishable. For example, in genetic study of the color of the hull in oat seeds, a variety with white hulls was crossed with a variety with black hulls. The F_1 hybrid seeds had black hulls. Among 560 progeny in the F_2 generation, the hull phenotypes observed were 418 black, 106 gray, and 36 white. The ratio of phenotypes is 11.6 : 3.9 : 1, or very nearly 12 : 3 : 1. A genetic hypothesis to explain these results is that the black-hull phenotype is due to the presence of a dominant allele A and the gray-hull phenotype is due to another dominant allele B whose effect is apparent only in aa genotypes. On the basis of this hypothesis, the original varieties had genotypes $aa\ bb$ (white) and $AA\ BB$ (black). The F_1 has genotype $Aa\ Bb$ (black). If the A, a allele pair and the B, b allele pair undergo independent assortment, then the F_2 generation is expected to have the genotypic and phenotypic composition 9/16 $A-\ B-$ (black hull), 3/16 $A-\ bb$ (black hull), 3/16 $aa\ B-$ (gray hull), 1/16 $aa\ bb$ (white hull), or 12 : 3 : 1.

13 : 3 is illustrated by the difference between White Leghorn chickens (genotype $CC\ Ii$) and White Wyandotte chickens (genotype $cc\ ii$). Both breeds have white feathers because the C allele is necessary for colored feathers, but the I allele in White Leghorns is a dominant inhibitor of feather coloration. The F_1 generation of a dihybrid cross between these breeds has the genotype $Cc\ Ii$, which is expressed as white feathers because of the inhibitory effects of the I allele. In the F_2 generation, only the $C-\ ii$ genotype has colored feathers, so there is a 13 : 3 ratio of white : colored.

9 : 4 : 3 is observed when homozygosity for a recessive allele with respect to one gene masks the expression of the genotype of a different gene. For example, if the aa genotype has the same phenotype regardless of whether the genotype is $B-$ or bb , then the 9 : 4 : 3 ratio results. As an example, in mice the grayish “agouti” coat color results from a horizontal band of yellow pigment just beneath the tip of each hair. The agouti pattern is due to a dominant allele A , and in aa animals the coat color is

black. A second dominant allele, C , is necessary for the formation of hair pigments of any kind, and cc animals are albino (white). In a cross of $AA\ CC$ (agouti) $\times$ $aa\ cc$ (albino), the F_1 progeny are $Aa\ Cc$ and phenotypically agouti. Crosses between F_1 males and females produce F_2 progeny in the proportions 9/16 $A-\ C-$ (agouti), 3/16 $A-\ cc$ (albino), 3/16 $aa\ C-$ (black), 1/16 $aa\ cc$ (albino), or 9 agouti : 4 albino : 3 black.

9 : 6 : 1 implies that homozygosity for either of two recessive alleles yields the same phenotype but that the phenotype of the double homozygote is different. In Duroc–Jersey pigs, red coat color requires the presence of two dominant alleles R and S . Pigs of genotype $R-\ ss$ and $rr\ S-$ have sandy-colored coats, and $rr\ ss$ pigs are white. The F_2 ratio is therefore 9/16 $R-\ S-$ (red), 3/16 $R-\ ss$ (sandy), 3/16 $rr\ S-$ (sandy), 1/16 $rr\ ss$ (white), or 9 red : 6 sandy : 1 white.

09131_01_1749P

3.7 Genetic Analysis: Mutant Screens and the Complementation Test

When a geneticist makes a statement such as “Flower color in peas is determined by a single gene with dominant allele P and recessive allele p ,” this does not mean that only one gene is needed for flower color. Implicit in the statement is the existence of two strains or varieties, one with colored flowers and one with white flowers, in which the difference is caused by one strain having genotype PP and the other pp . Many genes beside P are also necessary for purple flower coloration. Among them are other genes in the biochemical pathway for the synthesis of anthocyanin, as well as different genes needed for expression of the biochemical pathway in developing flowers. A third variety of peas that has white flowers may not have the genotype pp . In this variety, the white flowers could be caused by a mutation in one of the other genes needed for flower color. We have already seen an example in Figure 3.23, in which the white flower at the upper right

Photocaption photocaptionphotocaptionphoto-
captionphotocaptionphotocaptionphotocaption-
photocaptionphotocaptionphotocaptionphotoca-
ptionphotocaptionphotocaptionphotocaption-
photocaptionphotocaptionphotocaption

In a mutant screen for flower color, most of the new mutations will be recessive. This is because most of the new mutant alleles will encode an inactive protein of some kind, and most protein-coding genes are expressed at such a level that one copy of the wildtype allele is sufficient to yield a normal morphological phenotype. (The molecular phenotype is often intermediate; the heterozygotes have half as much protein as the wildtype homozygotes.) New recessive alleles are therefore identified in homozygous genotypes, because homozygous recessive plants have white flowers. Among all the new mutant strains that are isolated, some will have a recessive mutation in one gene, others in a second gene, still others in a third gene, and so on. Complicating the issue are multiple alleles, because two or more independently isolated mutations may be alleles of the same wildtype gene. How can the strains be sorted out? How can the geneticist determine which mutations are alleles of the same gene and which are mutations in different genes?

In part B we suppose the parental mutant strains are homozygous for recessive alleles of different genes, say a_1a_1 and b_1b_1 . With respect to the B locus, the genotype of the a_1a_1 strain is BB , and with respect to the A locus, the genotype of the b_1b_1 strain is

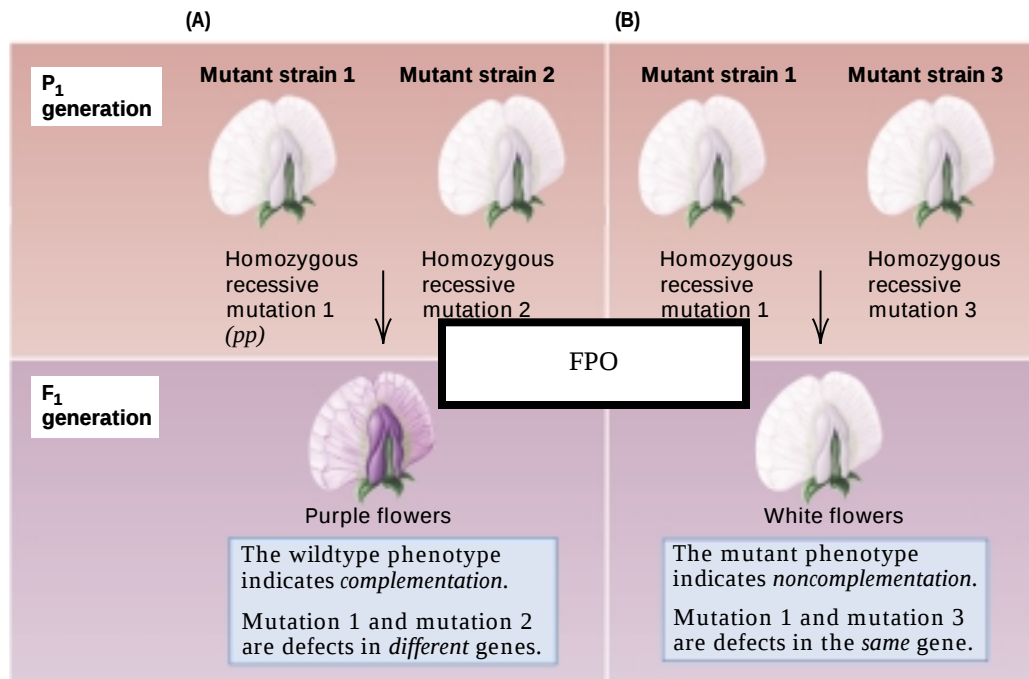


Figure 3.25 The complementation test reveals whether two recessive mutations are alleles of the same gene. In the complementation test, homozygous recessive genotypes are crossed. If the phenotype of the F_1 progeny is mutant (A), it means that the mutations in the parental strains are alleles of same gene. If the phenotype of the F_1 progeny is nonmutant (B), it means that the mutations in the parental strains are alleles of different genes.

AA; hence the genotypes of the mutant strains could be written more completely as $a_1a_1 BB$ and $AA b_1b_1$. In this case, the F_1 progeny have the genotype $Aa_1 Bb_1$. The A allele masks the mutation a_1 and the B allele masks the mutation b_1 ; hence the phenotype is purple (wildtype). This result is called **complementation**, and it means that the parental strains are homozygous for recessive alleles of *different* genes.

The kind of cross illustrated in Figure 3.25 is a **complementation test**. As we have seen, it is used to determine whether recessive mutations in each of two different strains are alleles of the same gene. Because the result indicates the presence or absence of allelism, the complementation test is one of the key experimental operations in genetics. To illustrate the application of the test in practice, suppose a mutant screen were carried out to isolate more mutations for white flowers. Starting with a true-breeding strain with purple flowers, we treat pollen with x rays and use the irradiated pollen to fertilize ovules to obtain seeds. The F_1 seeds are grown and

the resulting plants allowed to self-fertilize, after which the F_2 plants are grown. A few of the F_1 seeds may contain a new recessive mutation, but in this generation the genotype is heterozygous. Self-fertilization of such plants results in F_2 progeny in a ratio of 3 purple : 1 white. Because mutations resulting in a particular phenotype are quite rare, even when induced by radiation, only a few among many thousands of self-fertilized plants will be found to have a new white-flower mutation. Let us suppose that we were lucky enough to obtain six new mutations.

How are we going to name these mutations? We can make no assumptions about the number of genes represented. As far as we know at this stage, they could all be alleles of the same gene, or they could all be alleles of different genes. This is what we need to establish using the complementation test. For the moment, then, let us assign the mutations arbitrary names in series— $x1$, $x2$, $x3$, and so forth—with no implication about which may be alleles of each other.

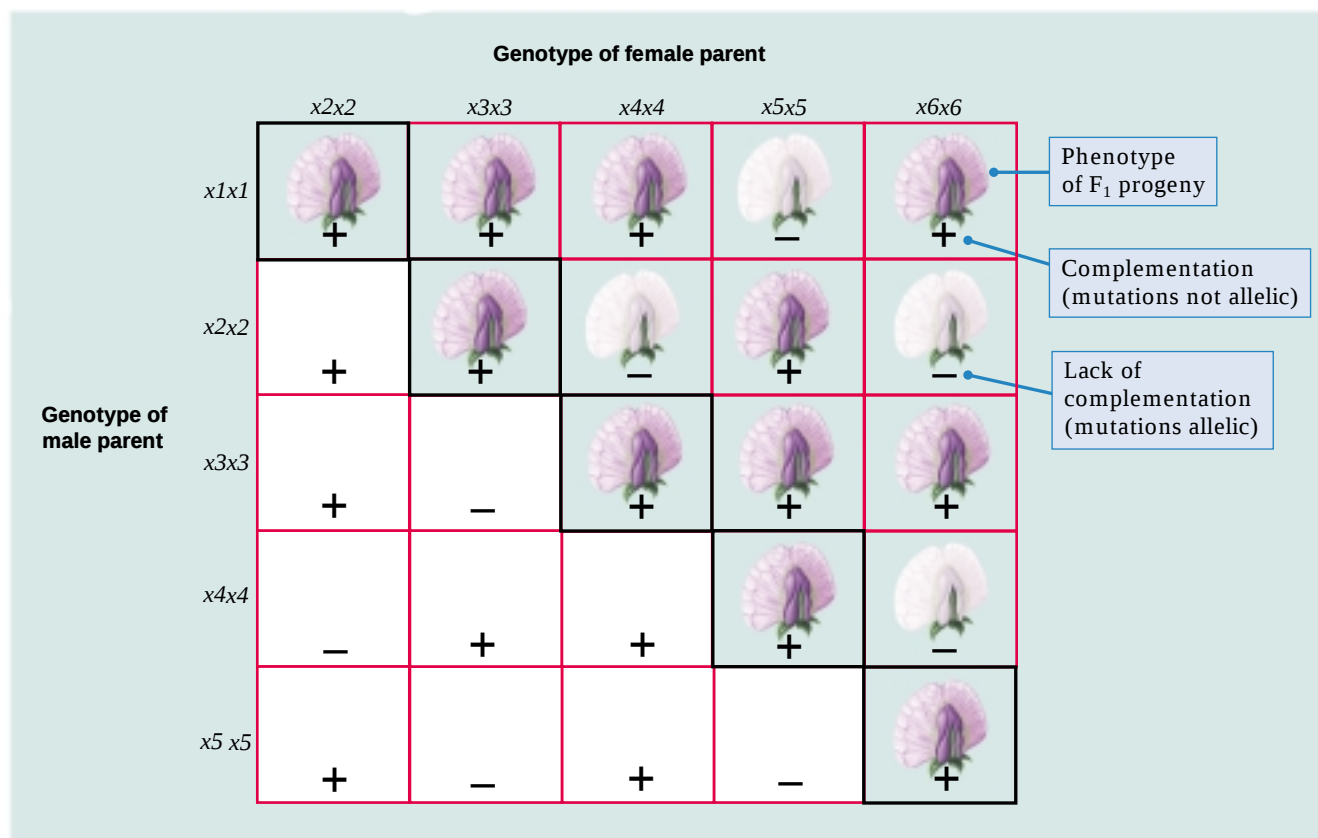


Figure 3.26 Results of complementation tests among six mutant strains of peas, each homozygous for a recessive allele resulting in white flowers. Each box gives the phenotype of the F₁ progeny of a cross between the male parent whose genotype is indicated in the far left column and the female parent whose genotype is indicated in the top row.

The next step is to classify the mutations into groups using the complementation test. **Figure 3.26** shows the results, conventionally reported in a triangular array of + and - signs. The crosses that yield F₁ progeny with the wildtype phenotype (in this case, purple flowers) are denoted with a + sign in the corresponding box, whereas those that yield F₁ progeny with the mutant phenotype (white flowers) are denoted with a - sign. The + signs indicate complementation between the mutant alleles in the parents, and the - signs indicate noncomplementation. (The bottom half of the triangle is unnecessary because reciprocal crosses yield the same results, and the diagonal elements are unnecessary because each strain is true-breeding within itself and so yields mutant progeny.) As we saw in **Figure 3.25**, complementation in a cross means that the parental strains have mutations in different genes. Lack of complementation means that the parental mutations are in the same

gene. The following principle underlies the complementation test.

The Principle of Complementation:

If two recessive mutations are alleles of the same gene, then a cross between homozygous strains yields F₁ progeny that are mutant (noncomplementation); if they are alleles of different genes, then the F₁ progeny are wildtype (complementation).

In interpreting complementation data such as those in **Figure 3.26**, we actually apply the principle the other way around. We examine the phenotype of the F₁ progeny of each possible cross and then infer whether or not the parental strains have mutant alleles of the same gene.

In a complementation test, if the combination of two recessive mutations results in a mutant phenotype, then the mutations are regarded as alleles of the same gene; if the combination results in a wildtype phenotype, then the mutations are regarded as alleles of different genes.

A convenient way to analyze the data in Figure 3.26 is to arrange the alleles in a circle as shown in Figure 3.27A. Then, for each possible pair of mutations, connect the pair by a straight line if the mutations *fail* to complement (Figure 3.27B). According to the principle of complementation, the lines must connect mutations that are alleles of each other, because in a complementation test, lack of complementation means that the mutations are alleles. Each of the groups of noncomplementing mutations is called a **complementation group**. As we have seen, each complementation group defines a gene, so the complementation test provides the geneticist's technical definition:

A gene is defined experimentally as a set of mutations that make up a single complementation group. Any pair of mutations within the group fail to complement one another, and the complementation test yields organisms with a mutant phenotype.

The mutations in Figure 3.27 therefore represent three genes, mutation of any one of which results in white flowers. Mutations *x1* and *x5* are alleles of one gene; *x3* is an allele of a different gene; and *x2*, *x4*, and *x6* are alleles of a third gene.

What about the genes *P* and *C* (Figure 3.23 on page 115) that also affect flower color? The complementation test tells us nothing about these, because we have not yet included them in any of the crosses. To test allelism with *P* and *C*, we would have to cross one mutant strain from each complementation group with strains of genotype *pp* and *cc*. If the *F*₁ progeny are mutant, it identifies the complementation group with an already known mutation. For example, suppose we cross each of the strains *x1x1*, *x3x3*, and *x4x4* with *pp* and *cc*. Suppose further that the progeny are all wildtype except in two crosses—*x1x1* × *pp* and *x3x3* × *cc*—in which the progeny are mutant. This result implies that the complementation group consisting of *x1* and *x5* also includes *p* and that the complementation group consisting of *x3* also includes *c*.

At this point in the genetic analysis, it is advisable to rename the mutations to indicate which ones are true alleles. (There is an old Chinese saying that the correct naming of things is the beginning of wisdom, and this is certainly true in the case of genes.) Because the *p* allele already had its

name before the mutation screen was carried out, the new alleles of *p* (that is, *x1* and *x5*) should be renamed to reflect their allelism with *p*. Old names have priority, and so we must use the symbol *p* embellished with some sort of identifier. We might, for example, rename the *x1* and *x5* mutations *p*₂ and *p*₃ to indicate that they were the second and third mutations of *P* to be discovered and to convey their independent origins. Similarly, we might rename *x3* (the new allele of *c*) as *c*₂. The remaining complementation group consisting of *x2*, *x4*, and *x6* is a new one, and we can name it arbitrarily. For example, we might call the locus *albus* (Latin for “white”) with the gene symbol *alb* and call the alleles *alb*₁, *alb*₂, and *alb*₃. The wildtype dominant allele of *alb*, which is necessary for purple coloration, would then be represented as *Alb* or as *alb*⁺. The procedure of sorting new mutations into complementation groups and renaming them according to their allelism is an example of how geneticists identify genes and name alleles. Such renaming of alleles is the typical manner in which genetic terminology evolves as knowledge advances.

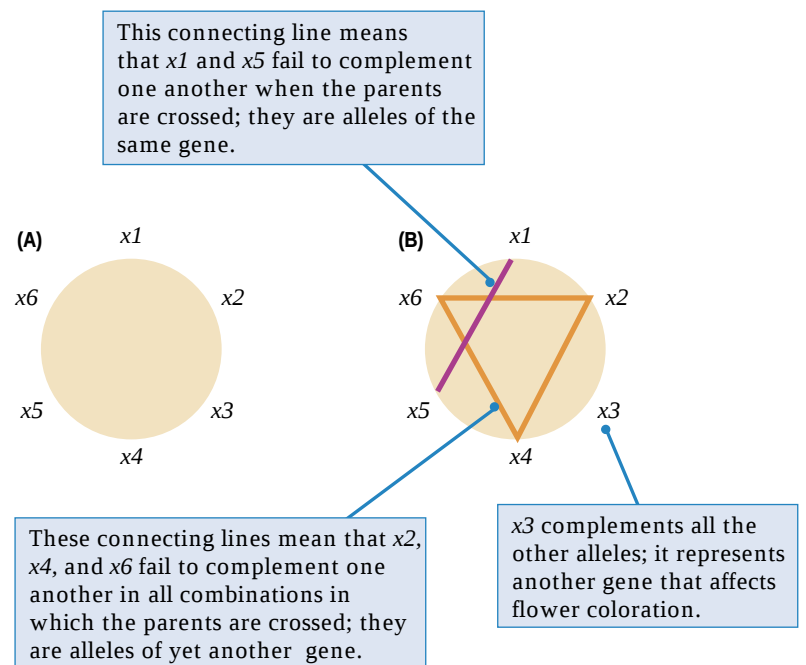


Figure 3.27 A method for interpreting the results of complementation tests. (A) Arrange the mutations in a circle. (B) Connect by a straight line any pair of mutations that fail to complement (that yield a mutant phenotype); any pair of mutations so connected are alleles of the same gene. In this example, there are three complementation groups, each of which represents a single gene needed for purple flower coloration.

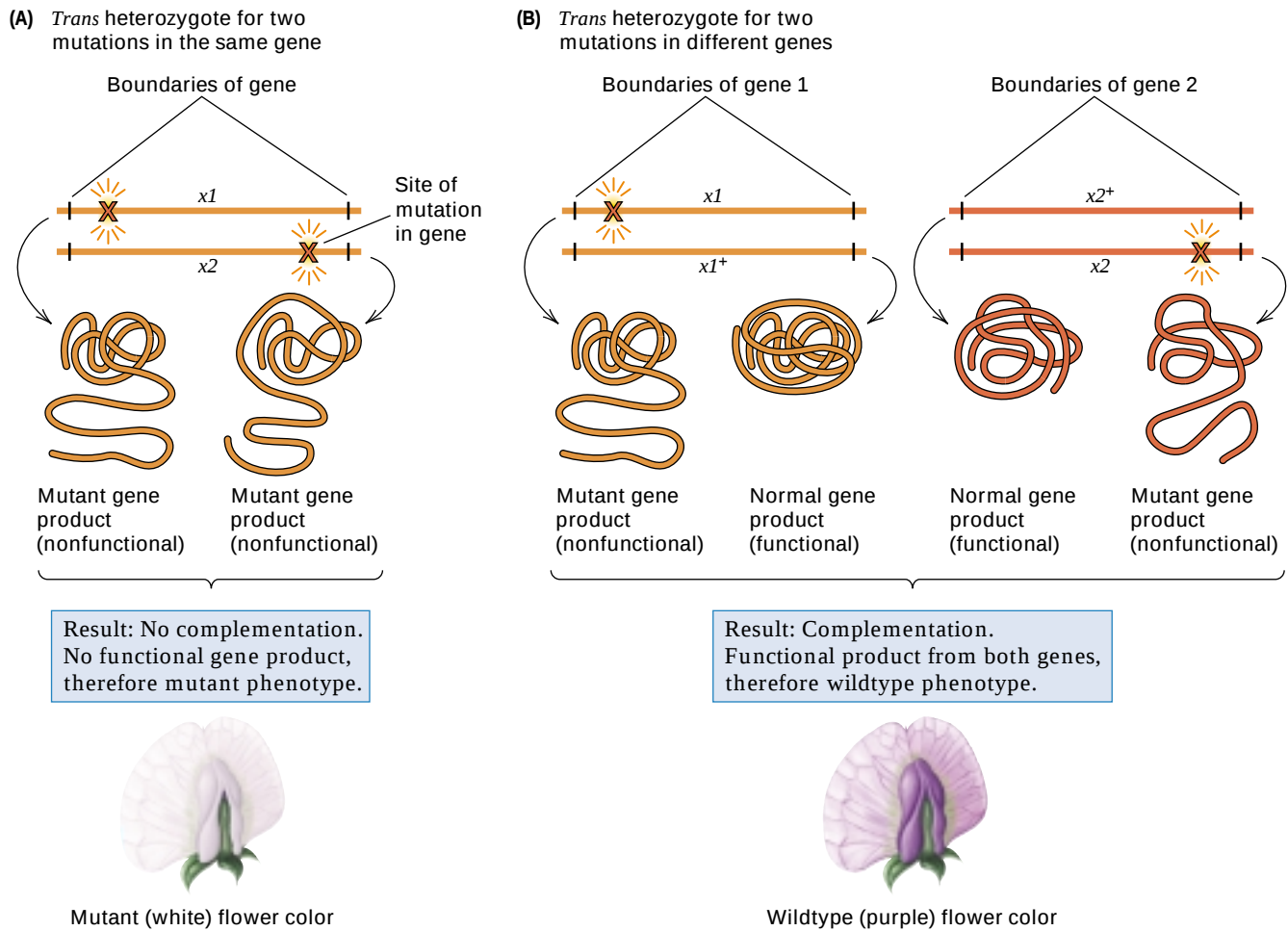


Figure 3.28 Molecular interpretation of a complementation test used in determining whether two mutations are alleles of the same gene (A) or alleles of different genes (B).

Why Does the Complementation Test Work?

The molecular basis of complementation is illustrated in **Figure 3.28**. Part A depicts the situation when two recessive mutations, $x1$ and $x2$, each resulting in white flowers, are different mutations in the same gene. (The site of each mutation is indicated by a “burst.”) The cross $x1x1 \times x2x2$ produces an F_1 hybrid in which $x1$ and $x2$ are in homologous DNA molecules; $x1$ encodes a protein with one type of defect, and $x2$ encodes a protein with a different type of defect, but both types of protein are nonfunctional. Hence alleles in the same gene yield a mutant phenotype (white flowers), because neither mutation encodes a wild-type form of the protein.

When the mutations are alleles of different genes, the situation is as depicted in **Figure 3.28B**. Because the mutations are in different genes, the homozygous $x1$ strain is also homozygous for the wildtype allele $x2^+$ of the second locus; likewise, the homozygous $x2$ strain is also homozygous for the wildtype allele $x1^+$ of the first locus. Hence, the same cross that yields the genotype $x1x2$ in the case of allelic mutations (**Figure 3.28A**) yields $x1^+x1\ x2^+x2$ in the case of different genes. This is a double heterozygous genotype. Because the mutations are both recessive and in different genes, they do complement each other and yield an organism with a wildtype phenotype (purple flowers). With respect to the protein rendered defective by $x1$, there is a functional form encoded by the

wildtype allele brought in from the $x2x2$ parent. With respect to the protein rendered defective by $x2$, there is again a functional form encoded by the wildtype

allele brought in from the $x1x1$ parent. Because a functional form of both proteins is produced, the result is a normal phenotype, or complementation.

Chapter Summary

Mendelian genetics deals with the hereditary transmission of genes from one generation to the next. One key principle is segregation, in which the two alleles in an individual separate during the formation of gametes so that each gamete is equally likely to contain either member of the pair. In a cross such as $AA \times aa$, in which only one gene is considered, the genotype of the offspring (constituting the F_1 generation) is heterozygous Aa . The phenotype of the F_1 progeny depends on the dominance relationships among the alleles. For many morphological traits, the wildtype allele, here denoted A , is dominant, and the phenotype of heterozygous Aa is indistinguishable from that of homozygous AA . In contrast, the alleles of molecular genetic markers are often codominant, and the phenotypes of AA , Aa , and aa are all distinct. For an RFLP marker, for example, the homozygous AA and aa genotypes have a phenotype consisting of a single band differing in electrophoretic mobility, whereas the heterozygous Aa genotype has a phenotype consisting of both bands.

In the formation of gametes, an Aa genotype produces A -bearing and a -bearing gametes in equal proportions. Hence in the $F_1 \times F_1$ cross $Aa \times Aa$, assuming random union of gametes in fertilization, the progeny (the F_2 generation) are expected to consist of genotypes $AA : Aa : aa$ in the proportions $1 : 2 : 1$. The distribution of phenotypes in the F_2 generation again depends on the dominance relationships. If A is dominant to a , then the F_2 ratio of dominant : recessive phenotypes is expected to be $3 : 1$. With codominance all three genotypes are distinguishable, and the ratio of F_2 phenotypes is $1 : 2 : 1$.

For crosses involving two genes—for example, $AA BB \times aa bb$ —the F_1 genotype is $Aa Bb$. Segregation of each gene implies that the ratios of $A : a$ and of $B : b$ gametes are both $1 : 1$. If the genes are unlinked, they undergo independent segregation (independent assortment), and the gametic types $AB : Ab : aB : ab$ are formed in the ratio $1 : 1 : 1 : 1$. Hence the F_2 generation formed from the cross $F_1 \times F_1$ is expected to have genotypes given by the product of the expression $(1/4 AA + 1/2 Aa + 1/4 aa) \times (1/4 BB + 1/2 Bb + 1/4 bb)$. Using a dash to represent an allele of unspecified type, we can write the F_2 genotypes as $9 A- B- : 3 A- bb : 3 aa B- : 1 aa bb$, and if both A and B are dominant, the phenotypic ratio in the F_2 generation is $9 : 3 : 3 : 1$. This ratio can be modified in various ways by interaction between the genes (epistasis). Different types of epistasis may result in dihybrid ratios such as $9 : 7$ or $12 : 3 : 1$ or $13 : 3$ or $9 : 4 : 3$.

The rules of probability provide the basis for predicting the outcomes of genetic crosses based on the principles of

segregation and independent assortment. Two basic rules for combining probabilities are the addition rule and the multiplication rule. The addition rule applies to mutually exclusive events; it states that the probability of the realization of either one or the other of two events equals the sum of the respective probabilities. The multiplication rule applies to independent events; it states that the probability of the simultaneous realization of both of two events is equal to the product of the respective probabilities.

In some organisms, including human beings, it is not possible to perform controlled crosses, and genetic analysis is accomplished through the study of pedigrees through two or more generations. Pedigree analysis is the determination of the possible genotypes of the family members in a pedigree and of the probability that an individual member has a particular genotype. The goal of pedigree analysis is often to infer the genetic basis of an inherited disease or other condition—for example, to determine whether it may be due to a simple dominant or recessive allele. Although most morphological traits do not show simple Mendelian patterns of inheritance in pedigrees, molecular markers usually do. Prominent among these are SNPs (single-nucleotide polymorphisms), RFLPs (restriction fragment length polymorphisms), and STRPs (simple tandem repeat polymorphisms).

Multiple alleles are often encountered in natural populations or as a result of mutant screens. Genes with multiple alleles may have as few as three, such as those for the ABO blood groups, to as many as a hundred or more. Examples of large numbers of alleles include the genes used in DNA typing and the self-sterility alleles in some flowering plants. Although there may be multiple alleles in a population, each gamete can carry only one allele of each gene, and each organism can carry at most two different alleles of each gene.

The complementation test is the functional definition of a gene. If a cross between two homozygous recessive genotypes results in nonmutant progeny, the mutations are said to *complement* one another. Complementation is evidence that the alleles are mutations in different genes. On the other hand, if a cross between two homozygous recessives results in mutant progeny, then the alleles show *noncomplementation* (they fail to complement). Noncomplementation is evidence that the mutations are alleles of the *same* gene. For any group of recessive mutations, a complete complementation test entails crossing the homozygous recessives in all pairwise combinations. At the molecular level, lack of complementation implies that allelic mutations impair the function of the same protein molecule.

Key Terms

addition rule	hybrid	segregation
antibody	incomplete dominance	sib
antigen	independent assortment	sibling
backcross	Mendelian genetics	testcross
carrier	multiple alleles	transmission genetics
codominance	multiplication rule	transposable element
complementation	mutant screen	true-breeding
complementation group	noncomplementation	wildtype
complementation test	outcrossing	unlinked genes
consanguineous mating	P ₁ generation	
epistasis	pedigree	
F ₁ generation	penetrance	
F ₂ generation	Punnett square	
gamete	reciprocal cross	

Review the Basics

- What is the principle of segregation, and how is this principle demonstrated in the results of a single-gene (monohybrid) cross?
- What is the principle of independent assortment, and how is this principle demonstrated in the results of a two-gene (dihybrid) cross?
- Explain why random union of male and female gametes is necessary for Mendelian segregation and independent assortment to be observed in the progeny of a cross.
- What is the difference between mutually exclusive events and independent events? How are the probabilities of these two types of events combined? Give two examples of genetic events that are mutually exclusive and two examples of genetic events that are independent.
- When two pairs of alleles show independent assortment, under what conditions will a 9 : 3 : 3 : 1 ratio of phenotypes in the F₂ generation *not* be observed?
- Explain the following statement: “Among the F₂ progeny of a dihybrid cross, the ratio of genotypes is 1 : 2 : 1, but among the progeny that express the dominant phenotype, the ratio of genotypes is 1 : 2.”
- What are the principal features of human pedigrees in which a rare dominant allele is segregating? In which a rare recessive allele is segregating?
- What is a mutant screen and how is it used in genetic analysis?
- Explain the statement: “In genetics, a gene is identified experimentally by a set of mutant alleles that fail to show complementation.” What is complementation? How does a complementation test enable a geneticist to determine whether two different mutations are or are not mutations in the same gene?
- What does it mean to say that epistasis results in a “modified dihybrid F₂ ratio?” Give two examples of a modified dihybrid F₂ ratio, and explain the gene interactions that result in the modified ratio.

Guide to Problem Solving

Problem 1 Explain what each of the following ratios represents in the F₂ progeny of a single-gene (monohybrid) cross. Which are ratios of genotypes and which are ratios of phenotypes?

- (a) 3 : 1
- (b) 1 : 2 : 1
- (c) 2 : 1

Answer (a) 3 : 1 is the expected ratio of phenotypes when one of the alleles is dominant. (b) 1 : 2 : 1 is the expected ratio of genotypes AA : Aa : aa. (c) 2 : 1 is the expected ra-

tio of heterozygous Aa genotypes to homozygous AA genotypes.

Problem 2 Complete the table by inserting 0, 1, 1/2, or 1/4 for the probability of each genotype of progeny from each type of mating. For which mating are the parents identical in genotype but the progeny as variable in genotype as they can be for a single locus? For which mating are the parents as different as they can be for a single locus but the progeny identical to each other and different from either parent?

Mating	Progeny genotypes		
	AA	Aa	aa
AA × AA			
AA × Aa			
AA × aa			
Aa × Aa			
Aa × aa			
aa × aa			

Answer This table is to transmission genetics what the multiplication table is to arithmetic. It is fundamental to being able to solve almost any type of quantitative problem in transmission genetics. It needs to be thoroughly understood—memorized, if necessary!

Mating	Progeny genotypes		
	AA	Aa	aa
AA × AA	1	0	0
AA × Aa	1/2	1/2	0
AA × aa	0	1	0
Aa × Aa	1/4	1/2	1/4
Aa × aa	0	1/2	1/2
aa × aa	0	0	1

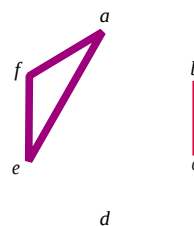
The mating for which the parents are identical in genotype, but the progeny are as variable in genotype as they can be for a single locus, is $Aa \times Aa$. The mating for which the parents are as different as they can be for a single locus, but the progeny are identical to each other and different from either parent, is $AA \times aa$.

Problem 3 The data in the accompanying complementation matrix summarize the result of crosses between mutants designated *a* through *f*.

- Configure the allele symbols in the form of a circle, and use straight lines to connect the alleles that are in the same complementation group.
- State the conclusion in words: How many complementation groups are indicated, and which mutants are in each complementation group?

	<i>a</i>	<i>b</i>	<i>c</i>	<i>d</i>	<i>e</i>	<i>f</i>
<i>a</i>	–	+	+	+	–	–
<i>b</i>		–	–	+	+	+
<i>c</i>			–	+	+	+
<i>d</i>				–	+	+
<i>e</i>					–	–
<i>f</i>						–

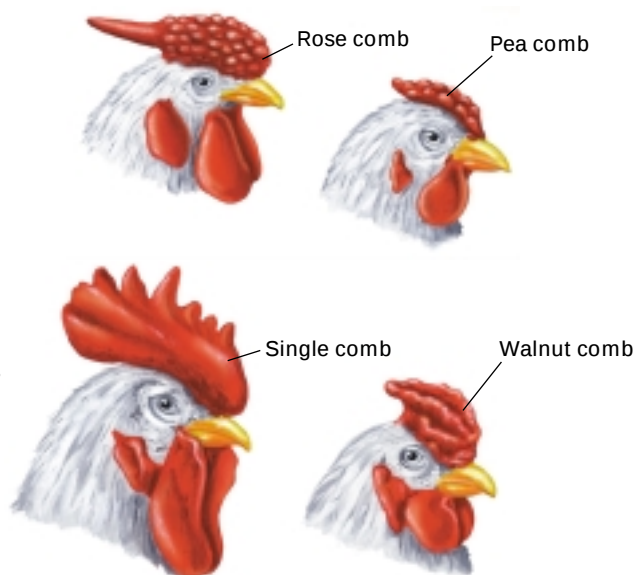
Answer (a) Follow the strategy outlined in the text, in which the alleles are placed in the shape of a circle and pairs of alleles that fail to show complementation are connected by a line. The resulting pattern is a visual representation of the complementation groups.



- The mutants define three complementation groups (“genes”). One complementation group consists of the alleles *a*, *e*, and *f*; another of the alleles *b* and *c*; and the other of the allele *d* only.

Problem 4 The accompanying illustration shows four alternative types of combs in chickens; they are called rose, pea, single, and walnut. The following data summarize the results of crosses. The rose and pea strains used in crosses 1, 2, and 5 are true breeding.

- rose × single → rose
- pea × single → pea
- (rose × single) $F_1 \times$ (rose × single) $F_1 \rightarrow$
3 rose : 1 single
- (pea × single) $F_1 \times$ (pea × single) $F_1 \rightarrow$
3 pea : 1 single
- rose × pea → walnut
- (rose × pea) $F_1 \times$ (rose × pea) $F_1 \rightarrow$
9 walnut : 3 rose : 3 pea : 1 single



- What genetic hypothesis can explain these results?
- What are the genotypes of parents and progeny in each of the crosses?
- What are the genotypes of true-breeding strains of rose, pea, single, and walnut?

Answer (a) Cross 6 gives the Mendelian ratios expected when two genes are segregating, so a genetic hypothesis with two genes is necessary. Crosses 1 and 3 give the results expected if rose comb were due to a dominant allele (say, *R*). Crosses 2 and 4 give the results expected if pea comb were due to a dominant allele (say, *P*). Cross 5 indicates that walnut comb results from the interaction of *R* and *P*. The segregation in cross 6 means that *R* and *P* are not alleles of the same gene.

- (b) 1. $RR\ pp \times rr\ pp \rightarrow Rr\ pp$
 2. $rr\ PP \times rr\ pp \rightarrow rr\ Pp$
 3. $Rr\ pp \times Rr\ pp \rightarrow 3/4\ R-\ pp : 1/4\ rr\ pp$

4. $rr\ Pp \times rr\ Pp \rightarrow 3/4\ rr\ P- : 1/4\ rr\ pp$

5. $RR\ pp \times rr\ PP \rightarrow Rr\ Pp$

6. $Rr\ Pp \times Rr\ Pp \rightarrow$

$9/16\ R-\ P- : 3/16\ R-\ pp : 3/16\ rr\ P- : 1/16\ rr\ pp$

(c) The true-breeding genotypes are *RR pp* (rose), *rr PP* (pea), *rr pp* (single), and *RR PP* (walnut).

Au: Should genotype in prob. 3.2 be "*Aa Bb . . .*" instead of "*AA Bb . . .*"?

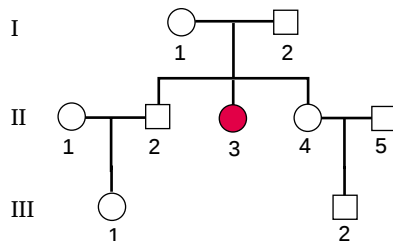
Analysis and Applications

3.1 What gametes can be formed by an individual organism of genotype *Aa*? Of genotype *Bb*? Of genotype *Aa Bb*?

3.2 How many different gametes can be formed by an organism with genotype *AA Bb Cc Dd Ee* and, in general, by an organism that is heterozygous for *m* genes and homozygous for *n* genes?

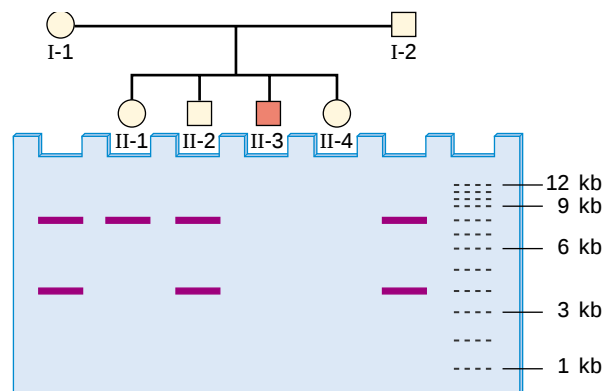
3.3 Round pea seeds are planted that were obtained from the F_2 generation of a cross between a true-breeding strain with round seeds and a true-breeding strain with wrinkled seeds. The pollen was collected and used *en masse* to fertilize plants from the true-breeding wrinkled strain. What fraction of the progeny is expected to have wrinkled seeds?

3.4 Phenylketonuria is a recessive inborn error of metabolism of the amino acid phenylalanine that results in severe mental retardation of affected children. The female II-3 (red circle) in the pedigree shown here is affected. If persons III-1 and III-2 (they are first cousins) mate, what is the probability that their offspring will be affected? (Assume that persons II-1 and II-5 are homozygous for the normal allele.)



3.5 Shown above right are a pedigree and gel diagram indicating the clinical phenotypes with respect to phenylketonuria and the molecular phenotypes with respect to an RFLP that overlaps the *PAH* gene for phenylalanine hydroxylase. Individual II-3 is affected.

- (a) Indicate the expected molecular phenotype of II-3.
 (b) Indicate the possible molecular phenotypes of II-4.



3.6 Assuming equal numbers of boys and girls, if a mating has already produced a girl, what is the probability that the next child will be a boy? If a mating has already produced two girls, what is the probability that the next child will be a boy? On what type of probability argument do you base your answers?

3.7 Assuming equal sex ratios, what is the probability that a sibship of four children consists entirely of boys? Of all boys or all girls? Of equal numbers of boys and girls?

3.8 In the following questions, you are asked to deduce the genotype of certain parents in a pedigree. The phenotypes are determined by dominant and recessive alleles of a single gene.

- (a) A homozygous recessive results from the mating of a heterozygote and a parent with the dominant phenotype? What does this tell you about the genotype of the parent with the dominant phenotype?
 (b) Two parents with the dominant phenotype produce nine offspring. Two have the recessive phenotype. What does this tell you about the genotype of the parents?
 (c) One parent has a dominant phenotype, and the other has a recessive phenotype. Two offspring result, and both have the dominant phenotype. What genotypes are possible for the parent with the dominant phenotype?

GeNETics on the Web will introduce you to some of the most important sites for finding genetic information on the Internet. To explore these sites, visit the Jones and Bartlett home page at

<http://www.jbpub.com/genetics>

For the book *Genetics: Analysis of Genes and Genomes*, choose the link that says **Enter GeNETics on the Web**. You will be presented with a chapter-by-chapter list of highlighted keywords. Select any highlighted keyword and you will be linked to a Web site containing genetic information related to the keyword.

- Mendel's paper is one of the few nineteenth century scientific papers that reads almost as clearly as if it had been written today. It is important reading for every aspiring geneticist. You can access a conveniently annotated text using the keyword **Mendel**. Although modern geneticists make a clear distinction between genotype and phenotype, Mendel made no clear distinction between these concepts. At this keyword site you will find a treasure trove of information about Mendel, including his famous paper, essays, commentary, and a collection of images—all richly linked to additional Internet resources.

- **Huntington disease** is a devastating degeneration of the brain that begins in middle life. It affects about 30,000 Americans, and, because of its dominant mode of inheritance and complete penetrance, each of their 150,000 siblings and children has a 50–50 chance of developing the disease. Named after George Huntington, a Long Island physician who first described it in 1872, the disease's principal symptom is an involuntary, jerky motion of the head, trunk, and limbs called *chorea*, after the Greek word for “dance.” At this keyword site

you can learn about genetic testing programs, early symptoms, and the time course of the disease.

- The red and purple colors of flowers, as well as of autumn leaves, result from members of a class of pigments called anthocyanins. The biochemical pathway for **anthocyanin** synthesis in the snapdragon, *Antirrhinum majus*, can be found at this keyword site. The enzyme responsible for the first step in the pathway (chalcone synthase) limits the amount of pigment formed, which explains why red and white flowers in *Antirrhinum* show incomplete dominance.

- The Mutable Site changes frequently. Each new update includes a different site that highlights genetics resources available on the World Wide Web. Select the **Mutable Site** for Chapter 3 and you will be linked automatically.

- The Pic Site showcases some of the most visually appealing genetics sites on the World Wide Web. To visit the genetics Web site pictured below, select the **PIC Site** for Chapter 3.



3.9 Pedigree analysis tells you that a particular parent may have the genotype $AA BB$ or $AA Bb$, each with the same probability. Assuming independent assortment, what is the probability of this parent's producing an Ab gamete? What is the probability of the parent's producing an AB gamete?

3.10 Assume that the trihybrid cross $AA BB rr \times aa bb RR$ is made in a plant species in which A and B are dominant but there is no dominance between R and r . Consider the F_2 progeny from this cross, and assume independent assortment.

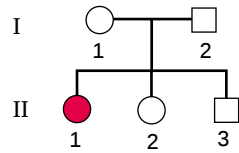
- How many phenotypic classes are expected?
- What is the probability of the parental $aa bb RR$ genotype?
- What proportion would be expected to be homozygous for all three genes?

3.11 In the cross $Aa Bb Cc Dd \times Aa Bb Cc Dd$, in which all genes undergo independent assortment, what proportion of offspring are expected to be heterozygous for all four genes?

3.12 The pattern of coat coloration in dogs is determined by the alleles of a single gene, with S (solid) being dominant over s (spotted). Black coat color is determined by the dominant allele A of a second gene, tan by homozygosity for the recessive allele a . A female having a solid tan coat is mated with a male having a solid black coat and produces a litter of six pups. The phenotypes of the pups are 2 solid tan, 2 solid black, 1 spotted tan, and 1 spotted black. What are the genotypes of the parents?

3.13 In the human pedigree shown here, the daughter indicated by the red circle (II-1) has a form of deafness determined by a recessive allele. What is the probability

that the phenotypically normal son (II-3) is heterozygous for the gene?

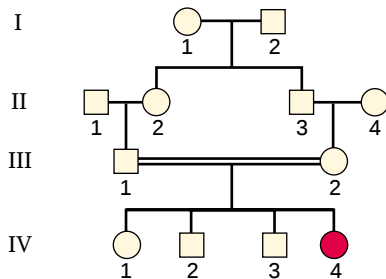


3.14 Huntington disease is a rare neurodegenerative human disease determined by a dominant allele, *HD*. The disorder is usually manifested after the age of forty-five. A young man has learned that his father has developed the disease.

- What is the probability that the young man will later develop the disorder?
- What is the probability that a child of the young man carries the *HD* allele?

3.15 Assume that the trait in the accompanying pedigree is due to simple Mendelian inheritance.

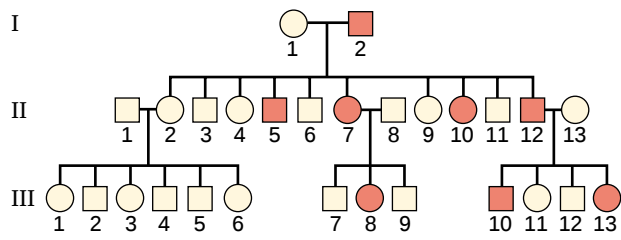
- Is it likely to be due to a dominant allele or a recessive allele? Explain.
- What is the meaning of the double horizontal line connecting III-1 with III-2?
- What is the biological relationship between III-1 and III-2?
- If the allele responsible for the condition is rare, what are the most likely genotypes of all of the persons in the pedigree in generations I, II, and III? (Use *A* and *a* for the dominant and recessive alleles, respectively.)



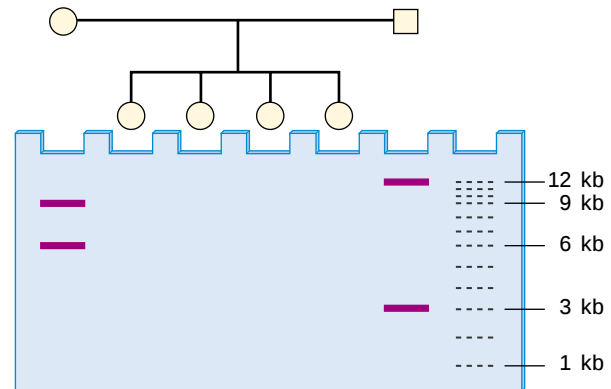
3.16 The Hopi, Zuni, and some other Southwest American Indians have a relatively high frequency of albinism (absence of skin pigment) resulting from homozygosity for a recessive allele, *a*. A normally pigmented man and woman, each of whom has an albino parent, have two children. What is the probability that both children are albino? What is the probability that at least one of the children is albino?

3.17 Say the trait in the accompanying pedigree is due to simple Mendelian inheritance.

- Is it likely to be due to a dominant allele or a recessive allele? Explain.
- What are the most likely genotypes of all of the persons in the pedigree? (Use *A* and *a* for the dominant and recessive alleles.)



3.18 The accompanying pedigree and gel diagram show the phenotypes of the parents for an RFLP that has multiple alleles. What are the possible phenotypes of the progeny, and in what proportions are they expected?



3.19 Red kernel color in wheat results from the presence of at least one dominant allele of each of two independently segregating genes (in other words, *R- B-* genotypes have red kernels). Kernels on *rr bb* plants are white, and the genotypes *R- bb* and *rr B-* result in brown kernel color. Suppose that plants of a variety that is true breeding for red kernels are crossed with plants true breeding for white kernels.

- What is the expected phenotype of the F_1 plants?
- What are the expected phenotypic classes in the F_2 progeny and their relative proportions?

3.20 Heterozygous *Cp cp* chickens express a condition called creeper, in which the leg and wing bones are shorter than normal (*cp cp*). The dominant *Cp* allele is lethal when homozygous. Two alleles of an independently segregating gene determine white (*W-*) versus yellow (*ww*) skin color. From matings between chickens heterozygous for both of these genes, what phenotypic classes will be represented among the viable progeny, and what are their expected relative frequencies?

3.21 White Leghorn chickens are homozygous for a dominant allele, *C*, of a gene responsible for colored feathers, and also for a dominant allele, *I*, of an independently segregating gene that prevents the expression of *C*. The White Wyandotte breed is homozygous recessive for both genes *cc ii*. What proportion of the F_2 progeny obtained from mating White Leghorn $\times$ White Wyandotte F_1 hybrids would be expected to have colored feathers?

3.22 The F_2 progeny from a particular cross exhibit a modified dihybrid ratio of 9 : 7 (instead of 9 : 3 : 3 : 1). What phenotypic ratio would be expected from a test-cross of the F_1 ?

3.23 Black hair in rabbits is determined by a dominant allele, B , and white hair by homozygosity for a recessive allele, b . Two heterozygotes mate and produce a litter of three offspring.

- What is the probability that the offspring are born in the order white-black-white? What is the probability that the offspring are born in either the order white-black-white or the order black-white-black?
- What is the probability that exactly two of the three offspring will be white?

3.24 Consider sibships consisting of 6 children, and assume a sex ratio of 1 : 1.

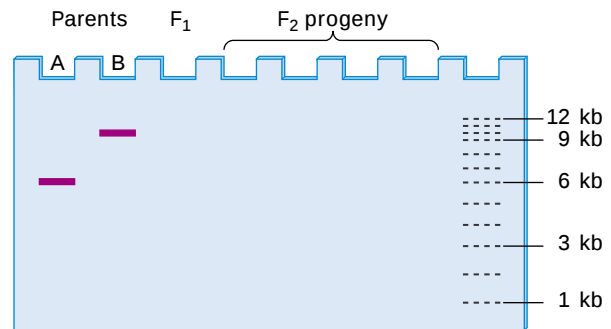
- What is the proportion with no girls?
- What is the proportion with exactly 1 girl?
- What is the proportion with exactly 2 girls?
- What is the proportion with exactly 3 girls?
- What is the proportion with 3 or more boys?

3.25 Andalusian fowls are colored black, splashed white (resulting from an uneven sprinkling of black pigment through the feathers), or slate blue. Black and splashed white are true breeding, and slate blue is a hybrid that segregates in the ratio 1 black : 2 slate blue : 1 splashed white. If a pair of blue Andalusians is mated and the hen lays three eggs, what is the probability that the chicks hatched from these eggs will be one black, one blue, and one splashed white?

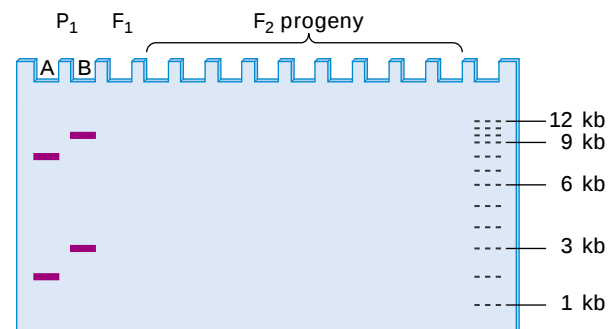
3.26 In the mating $Aa \times Aa$, what is the smallest number of offspring, n , for which the probability of at least one aa offspring exceeds 95 percent?

3.27 From the F_2 generation of a cross between mouse genotypes $AA \times aa$, one male progeny of genotype $A-$ was chosen and mated with an aa female. All of the progeny in the resulting litter were $A-$. From this result you would like to conclude that the sire's genotype is AA . How much confidence could you have in this conclusion for each litter size from 1 to 15? (In other words what is the probability that the sire's genotype is AA , given that the a priori probability is $1/3$ and that a litter of n pups resulted in all $A-$ progeny?)

3.28 The accompanying gel diagram includes the phenotype of two parents (A and B) with respect to two different RAPD polymorphisms. Each parent is homozygous for an allele associated with a band defining a RAPD polymorphism. The two RAPD polymorphisms are at different loci and undergo independent assortment. In the gel diagram, indicate the expected phenotype of the F_1 progeny as well as all possible phenotypes of the F_2 progeny, along with their expected proportions.



3.29 The gel diagram shown below shows the phenotype of two parents (A and B), each homozygous for two RFLPs that undergo independent assortment. Parent A has genotype $A_1A_1 B_1B_1$, where the A_1 allele yields a band of 2 kb and the B_1 allele yields a band of 8 kb. Parent B has genotype $A_2A_2 B_2B_2$, where the A_2 allele yields a band of 3 kb and the B_2 allele yields a band of 10 kb. Show the expected phenotype of the F_1 progeny as well as all possible phenotypes of the F_2 progeny, along with their expected proportions.

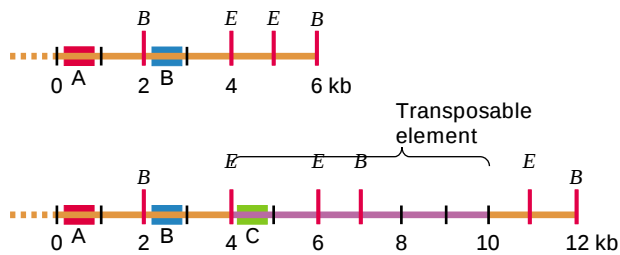


3.30 Complementation tests of the recessive mutant genes a through f produced the data in the accompanying matrix. The circles represent missing data. Assuming that all of the missing mutant combinations would yield data consistent with the entries that are known, complete the table by filling each circle with a + or - as needed.

	a	b	c	d	e	f
a	○	+	-	○	+	○
b	○	○	○	○	○	-
c	○	○	○	-	○	○
d	○	○	○	○	○	○
e	○	○	○	○	○	+
f	○	○	○	○	○	○

Challenge Problems

331 Diagrammed here is DNA from a wildtype gene (top) and a mutant allele (bottom) that has an insertion of a transposable element that inactivates the gene. The transposable element is present in many copies scattered throughout the genome. The symbols *B* and *E* represent the positions of restriction sites for *Bam*HI and *Eco*RI, respectively, and the rectangles show sites of hybridization with each of three probes (A, B, and C) that are available. The dots at the left indicate that the nearest site of either *Bam*HI or *Eco*RI cleavage is very far to the left of the region shown. Explain which probe and which single restriction enzyme you would use for RFLP analysis to identify both alleles. Also explain why any other choices would be unsuitable.



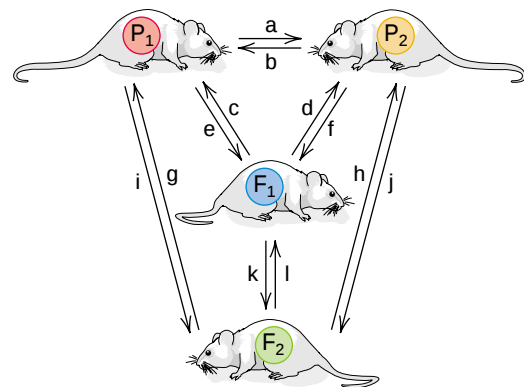
332 Meiotic drive is an unusual phenomenon in which two alleles do not show Mendelian segregation from the heterozygous genotype. Examples are known from mammals, insects, fungi, and other organisms. The usual mechanism is one in which both types of gametes are formed, but one of them fails to function normally. The excess of the driving allele over the other can range from a small amount to nearly 100 percent. Suppose that *D* is an allele showing meiotic drive against its alternative allele *d*, and suppose that *Dd* heterozygotes produce functional *D*-bearing and *d*-bearing gametes in the proportions 3/4 : 1/4. In the mating *Dd* × *Dd*,

- What are the expected proportions of *DD*, *Dd*, and *dd* genotypes?
- If *D* is dominant, what are the expected proportions of *D*- and *dd* phenotypes?
- Among the *D*- phenotypes, what is the ratio of *DD* : *Dd*?
- Answer parts (a) through (c), assuming that the meiotic drive takes place in only one sex.

333 The accompanying table summarizes the effect of inherited tissue antigens on the acceptance or rejection of transplanted tissues, such as skin grafts, in mammals. The tissue antigens are determined in a codominant fashion, so that tissue taken from a donor of genotype *Aa* carries both the *A* and the *a* antigen. In the table, the + sign means that a graft of donor tissue is accepted by the recipient, and the – sign means that a graft of donor tissue is rejected by the recipient. The rule is that *any graft will be rejected whenever the donor tissue contains an antigen not present in the recipient*. In other words, any transplant will be accepted if, and only if, the donor tissue does not contain an antigen different from any already present in the recipient.

		Donor		
		<i>AA</i>	<i>Aa</i>	<i>aa</i>
Recipient	<i>AA</i>	+	–	–
	<i>Aa</i>	+	+	+
	<i>aa</i>	–	–	+

The diagram illustrated below shows all possible skin grafts between inbred (homozygous) strains of mice (*P*₁ and *P*₂) and their *F*₁ and *F*₂ progeny. Assume that the inbred lines *P*₁ and *P*₂ differ in only one tissue-compatibility gene. For each of the arrows, what is the probability of acceptance of a graft in which the donor is an animal chosen at random from the population shown at the base of the arrow and the recipient is an animal chosen at random from the population indicated by the arrowhead?



Further Reading

Bowler, P. J. 1989. *The Mendelian Revolution*. Baltimore, MD: Johns Hopkins University Press.

Carlson, E. A. 1987. *The Gene: A Critical History*. 2d ed. Philadelphia: Saunders.

Dunn, L. C. 1965. *A Short History of Genetics*. New York: McGraw-Hill.

Hartl, D. L., and V. Orel. 1992. What did Gregor Mendel think he discovered? *Genetics* 131: 245.

Hawley, R. S., and C. A. Mori. 1999. *The Human Genome: A User's Guide*. San Diego: Academic.

Henig, R. M. 2000. *The Monk in the Garden*. Boston, MA: Houghton Mifflin.

- Judson, H. F. 1996. *The Eighth Day of Creation: The Makers of the Revolution in Biology*. Cold Spring Harbor, NY: Cold Spring Harbor Laboratory Press.
- Lewis, R. 2001. *Human Genetics: Concepts and Applications*. 4th ed. Dubuque, IA: McGraw-Hill.
- Lewontin, R. C. 2000. *It Ain't Necessarily So: The Dream of the Human Genome and Other Illusions*. New York: New York Review of Books.
- Maroni, G. 2000. *Molecular and Genetic Analysis of Human Traits*. Malden, MA: Blackwell.
- Mendel, G. 1866. Experiments in plant hybridization. (Translation.) In *The Origins of Genetics: A Mendel Source Book*, ed. C. Stern and E. Sherwood. 1966. New York: Freeman.
- Olby, R. C. 1966. *Origins of Mendelism*. London: Constable.
- Orel, V. 1996. *Gregor Mendel: The First Geneticist*. Oxford, England: Oxford University Press.
- Orel, V., and D. L. Hartl. 1994. Controversies in the interpretation of Mendel's discovery. *History and Philosophy of the Life Sciences* 16: 423.
- Sandler, I. 2000. Development: Mendel's legacy to genetics. *Genetics* 154: 7.
- Stern, C., and E. Sherwood, eds. 1966. *The Origins of Genetics: A Mendel Source Book*. New York: Freeman.
- Sturtevant, A. H. 1965. *A Short History of Genetics*. New York: Harper & Row.
- Sykes, B, ed. 1999. *The Human Inheritance*. New York: Oxford University Press.